

Dynamics in stochastic evolutionary models

DAVID K. LEVINE

Department of Economics, Washington University in St. Louis and European University Institute

SALVATORE MODICA

Department of Economics, Università di Palermo Salvatore

We characterize transitions between stochastically stable states and relative ergodic probabilities in the theory of the evolution of conventions. We give an application to the fall of hegemonies in the evolutionary theory of institutions and conflict, and illustrate the theory with the fall of the Qing dynasty and the rise of communism in China.

KEYWORDS. Evolution, conventions, Markov chains, state power.

JEL CLASSIFICATION. C73.

1. INTRODUCTION

The modern theory of the evolution of conventions deals with a Markov process in which there are strong forces such as learning toward equilibrium and weaker evolutionary forces such as “mutations” that disturb an equilibrium and lead from one equilibrium to another. There is a large economics literature studying models of this type: just to take some recent contributions, [Kreindler and Young \(2012, 2013\)](#) and [Ellison et al. \(2014\)](#) show how convergence in these models may be very fast, and [Sabourian and Juang \(2012\)](#) give a folk theorem for equilibrium selection. To prove theorems, the limit when weak forces are small is analyzed. In the limit, equilibria appear as recurrent communicating classes of the Markov process: sets of states all of which are accessible to each other, but are grouped into classes that are isolated from each other. Prior to the limit—the situation of interest—the Markov process is ergodic and puts positive weight on all states. However some states are more equal than others, and in the unique limit of the stationary distributions weight is placed only on the recurrent communicating classes of the limit; moreover, only some of these classes have positive weight—the stochastically stable classes. In particular, while the limiting Markov process can have many classes, the limit of the Markov processes may place weight on only one or a few classes. The literature, especially [Foster and Young \(1990\)](#), [Kandori et al. \(1993\)](#), [Young](#)

David K. Levine: david@dklevine.com

Salvatore Modica: salvatore.modica@unipa.it

We are especially indebted to Juan Block for his many comments and suggestions. We would also like to thank Drew Fudenberg, Kevin Hasker, Matt Jackson, Peyton Young, and five anonymous referees. We are grateful to NSF Grant SES-08-51315 and to the MIUR PRIN 20103S5RN3 for financial support.

Copyright © 2016 David K. Levine and Salvatore Modica. Licensed under the [Creative Commons Attribution-NonCommercial License 3.0](#). Available at <http://econtheory.org>.

DOI: [10.3982/TE1978](https://doi.org/10.3982/TE1978)

(1993), Ellison (2000), Cui and Zhai (2010), and Hasker (2014), develops a set of techniques for determining which of these classes get weight in the limit and gives a useful picture of what the stochastic process looks like when the weak forces are small but not zero. Roughly, the classes that have positive weight in the limit are seen most of the time, but the system will occasionally move away from a class and back again or transit from one class to another.

The focus of much of the literature has been on determining which classes get weight in the limit, and this is important to understand, but the transitions—the movement from one class to another—are also interesting and important for economics. For example, in a model of evolution such as that of Levine and Modica (2013b), where different economic and political institutions compete with each other, the recurrent communicating classes correspond to hegemonies—a single society that controls all economic resources—and the stochastically stable classes are the most powerful hegemonies. There is considerable historical evidence for the existence of hegemonies: China, the Roman Empire, and so forth. In the theory, as in reality, hegemonies inevitably fall and eventually reappear. How the transitions take place, i.e., what kind of phases mark the crucial steps in the transitions, is of some interest. Do the eventual winners of the conflict appear on the scene and battle back and forth with the hegemony for awhile until they take over and establish their own hegemony (short answer, “no”) or does something else happen, and if so what? Our results make it possible to answer these questions in some detail, describing the rise and fall of hegemony and the warring states period that takes place in between. It may also serve as a guide to policy, showing how different institutions can impact transitions.

The mathematical methods used in analyzing stochastic stability contain clues for what the transitions might look like. In particular, stochastically stable classes can be characterized by trees of recurrent communicating classes, where the distance between classes is measured by “resistance” and the stochastically stable classes appear as the root of trees with least total resistance. Because of the role played by least resistant paths in this analysis, a natural conjecture is that least resistant paths are in some sense more likely than higher resistant paths. Here we establish in exactly what sense this is true.

The starting point is to observe that a basic feature of resistance is that if we compare the probability of two specific paths when evolutionary forces are very weak, the lower resistance path is far more likely than the higher resistance path. However, if we look at all paths from one recurrent communicating class to another—the “quasi-direct routes” that include those that may dawdle within the class they start from before moving on—then as a group least resistant paths are far *less* likely than higher resistant paths. The reason for this is simple: it is likely to take a very long time to reach another recurrent communicating class; in the meantime there are likely to be many failed attempts to get there, and these attempts will typically involve some resistance. However, consider the actual transition from one class to another, that is, the paths that leave the starting recurrent communicating class for the final time and that do not pass through a third

class: we call these *direct routes*.¹ We show that the transition to the target class is likely to happen relatively quickly, in the sense that the set of least resistance direct routes are far more likely than other paths.

We establish the theory in two parts: we first develop a set of bounds for direct routes and then for quasi-direct routes. As a by-product, understanding these transitions also gives us clues about the ergodic probabilities. By examining which recurrent communicating classes are reached “next” from a given starting point, we construct a straightforward recursive algorithm that gives precise bounds on the ratio between the ergodic probabilities of all states that are “reasonably close” to recurrent communicating classes.

To illustrate the theory, we apply it to a simplified version of the [Levine and Modica \(2013b\)](#) evolutionary model of conflict and the emergence of hegemonies—some details of which are motivated by the transition theory of [Acemoglu and Robinson \(2001\)](#)—and provide, in particular, an account of the fall of the last Qing dynasty in China and the ensuing rise of communism.

2. MAIN RESULTS THROUGH AN ILLUSTRATIVE EXAMPLE

We are interested in economic models that can be represented as Markov processes where some transitions are much less likely than others. To illustrate this, we start with an example of a “standard” evolutionary model. This simple and familiar example is designed to illustrate the gaps in current knowledge and how the results of this paper fill those gaps. A historical application that motivates why these gaps are interesting is examined in [Section 7](#). Consider the 2×2 symmetric coordination game with actions G , B and payoff matrix

	G	B
G	2, 2	0, 0
B	0, 0	1, 1

This game has two pure Nash equilibria at GG and BB , and a mixed equilibrium with probability $1/3$ of G .

To put this in an evolutionary context, we assume that there are five players. Each player receives a payoff equal to the average he gets in all matches against his four opponents.² In each period, a single player is chosen with equal probability to reconsider her move; the other four play as in the previous period. The state of the system is the number of players playing G . So the state space Z has $N = 6$ states. We first define the behavior rule representing “rational” learning:³ the player who gets to move

¹Note that “direct” is sometimes used to mean “in a single transition.” Here a direct route can pass through many transitions, but it must not pause in an ergodic class along the way. From [Ellison \(2000\)](#) we know that such paths are not necessarily the quickest way to get to the target, an issue we carefully account for.

²As [Ellison \(1993\)](#) points out, this global interaction model converges much more slowly than if each player is matched only with a neighbor. Here the model is intended for illustrative purposes.

³In some models such as [Kandori et al. \(1993\)](#), this rational component of the dynamic is deterministic so can be referred to as the *deterministic dynamic*. Here, as in [Binmore and Samuelson \(1997\)](#) and [Blume \(2003\)](#), the revision rule is stochastic because the player who moves is determined randomly.

chooses a best response to the actions chosen by the opposing players. In addition, there are independent trembles: with probability $1 - \epsilon$ the behavior rule is followed, while with probability ϵ the player's choice is uniform and random over all possible actions. The presumption is that the chance of "arational" play ϵ is small compared to the probability $1 - \epsilon$ of "rational" play. This arational play is often called a *mutation* in the literature.

This dynamic can be represented as a Markov process on the state space Z defined above with six states representing the number of players playing G . Denoting source states by rows and target states by columns as is standard in the theory of Markov chains, the transition matrix can be computed as⁴

$$P_\epsilon = \begin{pmatrix} 1 - \frac{1}{2}\epsilon & \frac{1}{2}\epsilon & 0 & 0 & 0 & 0 \\ (\frac{1}{5} - \frac{1}{10}\epsilon) & (\frac{4}{5} - \frac{3}{10}\epsilon) & \frac{4}{10}\epsilon & 0 & 0 & 0 \\ 0 & (\frac{2}{5} - \frac{2}{10}\epsilon) & \frac{1}{2}\epsilon & (\frac{3}{5} - \frac{3}{10}\epsilon) & 0 & 0 \\ 0 & 0 & \frac{3}{10}\epsilon & (\frac{3}{5} - \frac{1}{10}\epsilon) & (\frac{2}{5} - \frac{2}{10}\epsilon) & 0 \\ 0 & 0 & 0 & \frac{4}{10}\epsilon & (\frac{4}{5} - \frac{3}{10}\epsilon) & (\frac{1}{5} - \frac{1}{10}\epsilon) \\ 0 & 0 & 0 & 0 & \frac{1}{2}\epsilon & 1 - \frac{1}{2}\epsilon \end{pmatrix}.$$

A critical concept in analyzing this system for ϵ small but not 0 is the notion of resistance. Although we give a formal definition below, in this example the resistance is the minimum number of transitions with probability of order ϵ needed to get from one state to another. For example, the resistance of going from $\{1\}$ to $\{0\}$ is 0 since the transition probability is $P_\epsilon(1, 0) = \frac{1}{5} - \frac{1}{10}\epsilon$, while the resistance of going from $\{0\}$ to $\{1\}$ is 1 since the transition probability is $P_\epsilon(0, 1) = \frac{1}{2}\epsilon$. Transitions with probability zero independent of ϵ have infinite resistance. The resistance matrix in this case is, therefore,

$$r = \begin{pmatrix} 0 & 1 & \infty & \infty & \infty & \infty \\ 0 & 0 & 1 & \infty & \infty & \infty \\ \infty & 0 & 1 & 0 & \infty & \infty \\ \infty & \infty & 1 & 0 & 0 & \infty \\ \infty & \infty & \infty & 1 & 0 & 0 \\ \infty & \infty & \infty & \infty & 1 & 0 \end{pmatrix}.$$

We can analyze this model by giving a heuristic outline of the methods we will develop in the paper.

⁴Take, for example, the first diagonal entry $P_\epsilon(0, 0) = 1 - \frac{1}{2}\epsilon$: if the current state is 0, then whoever is picked next period faces four opponents playing B , so the best response is B . With probability $1 - \epsilon$, the player is rational and plays B ; with probability ϵ , the player is arational and plays B with probability $\frac{1}{2}$. So with probability $1 - \epsilon + \frac{1}{2}\epsilon = 1 - \frac{1}{2}\epsilon$, the move is B and the next state will be 0. As another example, take the second row: the only G player is drawn with probability $\frac{1}{5}$, while with probability $\frac{4}{5}$, it is one of the B players; the best response is still B in any case, so play probability is again $1 - \frac{1}{2}\epsilon$ for B and $\frac{1}{2}\epsilon$ for G ; hence the state remains at one G player if either the G player is chosen and plays G (probability $(\frac{1}{5}) \cdot (\frac{1}{2}\epsilon)$) or a B player is chosen and plays B (probability $(\frac{4}{5}) \cdot (1 - \frac{1}{2}\epsilon)$); the sum of these two terms is the $\frac{4}{5} - \frac{3}{10}\epsilon$ appearing at the second diagonal entry $P_\epsilon(1, 1)$ in the matrix.

1. Recurrent communicating classes for $\epsilon = 0$.⁵

The two recurrent communicating classes consist of the singleton sets $\{0\}$, $\{5\}$. These sets are each absorbing and here they correspond to the pure Nash equilibria of the game. We denote the set of recurrent communicating classes by $\Omega = \{\Omega_0, \Omega_5\} = \{\{0\}, \{5\}\}$.

2. Relation between the classes for $\epsilon > 0$.

Starting in one recurrent communicating class, which comes next, how long will it take and relatively how much time is spent at each recurrent communicating class?

Here there is only one other recurrent communicating class, so the next one is always the other one. In the general case, [Corollary 3](#) tells us that the “next one” is one that can be reached with least resistance. This least resistance is known as the radius. In this example, it takes two mutations to get from $\{0\}$ to $\{5\}$ and three to get back, so the radius of $\{0\}$ is 2 and the radius of $\{5\}$ is 3.

How long it will take to get from one class to another is known from [Ellison \(2000\)](#)—the main existing result concerning transitions—and is covered here in [Theorem 4](#): it is ϵ^{-1} raised to the power of the radius: for $\{0\}$, the waiting time to $\{5\}$ is of order ϵ^{-2} ; for $\{5\}$, the waiting time to $\{0\}$ is of order ϵ^{-3} . Relatively ϵ^{-1} times as long is spent at $\{5\}$ as at $\{0\}$. When there are many recurrent communicating classes, computing the relative amount of time at each is complicated: we give a constructive algorithm for finding it in [Section 6.3](#).

3. Basins.

The basin of a recurrent communicating class is the set of states for which the probability of eventually reaching that class when $\epsilon = 0$ is 1. The basin of $\{0\}$ consists of the points $\{0\}$, $\{1\}$. The basin of $\{5\}$ is $\{3\}$, $\{4\}$, $\{5\}$. The state $\{2\}$ is in the “outer range” of both $\{0\}$ and $\{5\}$, but the basin of neither: it has positive probability of reaching either of the two recurrent communicating classes.

4. Relation between basins and classes.

The basin consists of states that are “close” to the corresponding recurrent communicating class.⁶ By [Theorem 4](#), during the time before reaching a new class most of the time will be spent in the current recurrent communicating class, but there will be a large number of periods during which the system will move to these close points and back. From [Theorem 8](#), the relative amount of time spent at these states compared to the time spent in the recurrent communicating class is of the order of the difference between the resistance of reaching the point and the radius. For example, starting at $\{5\}$, the radius is 3 and the resistance of getting to $\{3\}$ from $\{5\}$ is 2 so that the system will spend roughly ϵ^{-1} times as much time at $\{5\}$ as at $\{3\}$. From

⁵A *closed* or *recurrent communicating class* of a Markov process is a set with the property that there is a positive probability of reaching any point in the set from any other and the probability of leaving the set is zero. The literature also sometimes refers to these as limit sets.

⁶Strictly speaking, this is true not of the basin, but only of the inner basin from [Definition 5](#). However, in this example, the inner basin and the basin coincide.

the result above, that also means that the system spends roughly the same amount of time at $\{3\}$ as at $\{0\}$, but the nature of the time is quite different. The system will remain at $\{0\}$ for long contiguous periods of time with only occasional departures, while the system will remain at $\{3\}$ for very short periods of time, but go there very frequently.

5. Transitions

How do we get from one recurrent communicating class to another? This is covered in [Theorem 3](#). It says that with very high probability, the path will have least “peak resistance”: such a path may leave the recurrent communicating class and return any number of times, but during each departure from the recurrent communicating class, the resistance encountered can be no more than the radius. When the recurrent communicating class is left for the final time, the path followed to the new recurrent communicating class must have least resistance and the transition is “very quick.” In the example, going from $\{5\}$ to $\{0\}$, the path can leave $\{5\}$ and return many times, but the resistance encountered during these departures cannot be greater than 3. So, for example, $(5, 4, 3, 2, 3, 4, 5)$ can occur since it has resistance 3, but $(5, 4, 3, 4, 3, 4, 3, 4, 5)$ cannot occur because it has resistance 4. The final transition must have resistance 3, which in this case means it must be monotone: $\{4\}, \{3\}, \{2\}, \{1\}, \{0\}$ must occur in that order. However, any of these states can recur except $\{2\}$, since remaining adds no resistance. So, for example, $(5, 4, 4, 3, 2, 1, 1, 0)$ is a possible transition, but $(5, 4, 3, 2, 2, 1, 0)$ is not. Despite the fact that the transition paths can have loops of zero resistance in them, the expected length of the path is bounded independent of ϵ . This is shown in [Theorem 1](#).

3. THE MODEL

In the general case, we are given a finite state space Z with N elements and a family P_ϵ of Markov chains⁷ on Z indexed by $0 \leq \epsilon < 1$. This family satisfies two regularity conditions:

1. We have $\lim_{\epsilon \rightarrow 0} P_\epsilon = P_0$.
2. For all $x, z \in Z$, there is a resistance function $0 \leq r(x, z) \leq \infty$ and constants $0 < C < 1 < D < \infty$ such that $C\epsilon^{r(x,z)} \leq P_\epsilon(z|x) \leq D\epsilon^{r(x,z)}$.

Notice that zero resistance is equivalent to positive probability with respect to P_0 —a fact we will use all the time—and infinite resistance is zero probability in all P_ϵ 's. If $f(\epsilon)$ and $g(\epsilon)$ are nonnegative functions, a useful notation concerning resistances is to define $f(\epsilon) \sim g(\epsilon)$ if $\liminf_{\epsilon \rightarrow 0} f(\epsilon)/g(\epsilon) > 0$ and $\limsup_{\epsilon \rightarrow 0} f(\epsilon)/g(\epsilon) < \infty$ with the obvious

⁷In the literature it is often assumed that for $\epsilon > 0$ the chain is ergodic: in our analysis of the limit ergodic distribution in the later part of the paper we make such an assumption. However, for the analysis of transitions the assumption is not needed and it can be useful to apply the analysis to interim dynamics at states that will never be reached again. Moreover, the assumption that the state space is finite is not needed for the analysis of transitions, and the bounds given in [Appendix A](#) hold for countable state spaces as well as finite state space; in this case, when $\epsilon > 0$, there may be no long-run limit at all.

convention that $0 \sim 0$. With this notation, we can then write $P_\epsilon(z|x) \sim \epsilon^{r(x,z)}$. For readability we will state results in the text using this order notation; we restate (and prove) the results with exact bounds in the [Appendices](#).

As in the example of the previous section, we let Ω be the set of the recurrent communicating classes of P_0 . We write Ω_x for the recurrent communicating class containing x , where $\Omega_x = \emptyset$ if x is not part of a recurrent communicating class. A path a is a finite sequence $(z_0, z_1, \dots, z_t)$ of at least two not necessarily distinct states in Z and we write $t(a) = t$: this is the number of transitions (z_{s-1}, z_s) . The resistance of the path is $r(a) \equiv r(z_0, z_1) + r(z_1, z_2) + \dots + r(z_{t-1}, z_t)$.

We summarize some well known properties of P_0 and Ω . Nonempty recurrent communicating classes $\Omega_x \neq \emptyset$ are characterized by the property that from any point $y \in \Omega_x$, there is a positive probability path to any other point $z \in \Omega_x$ and that every positive probability path starting at y must lie entirely within Ω_x . Since positive probability in P_0 is the same as zero resistance, we may equally say that from any point $y \in \Omega_x$ there is a zero-resistance path to any other point $z \in \Omega_x$ and that every zero-resistance path starting at $y \in \Omega_x$ must lie entirely within Ω_x . An additional useful notion is defined as follows.

DEFINITION 1. A set W is *comprehensive* if for any point $z \in Z$ there is a zero-resistance path $a = (z, \dots, w)$ to some point $w \in W$. In particular, the set Ω is comprehensive.

We can give the following characterization of a comprehensive set.

PROPOSITION 1. A set W is comprehensive if and only if it contains at least one point from every nonempty recurrent communicating class, that is, for all $\Omega_x \in \Omega$ there exists $w \in W$ with $w \in \Omega_x$.

PROOF. Sufficiency: For any point $z \in Z$, there is a zero-resistance path to some point y in some recurrent communicating class Ω_y , and from there a zero-resistance continuation to the point in $\Omega_y \cap W$ that is assumed to exist. Necessity: If there is a set $\Omega_y \neq \emptyset$ with $\Omega_y \cap W = \emptyset$, then the zero-resistance path to W assumption fails: any zero-resistance path originating in Ω_y must remain entirely within Ω_y and, hence, does not reach W . $\square$

Notice that there are a great many comprehensive sets and our analysis is conditional on a particular choice of a comprehensive set. Different comprehensive sets may serve different useful purposes.

4. DIRECT ROUTES

In P_0 , a path that hits a point in a recurrent communicating class is then trapped in that class, so cannot reach a target outside of that class. When $\epsilon > 0$, this need not be the case. However, if a point in a recurrent communicating class is hit, then it is very likely that the path will then linger in that recurrent communicating class, passing through every point in the class many times and, in particular, through states in any comprehensive set. Hence, there is a sense in which paths that do not hit a comprehensive set must be “quick”: they cannot linger in a recurrent communicating class. We will call such paths

direct routes. Now *any* path that leaves a class Ω_x contains a direct route—the route to its first point in $\Omega \setminus \Omega_x$. So we start by studying direct routes. They are a bit like the hare in the story of the tortoise and the hare. Direct routes get to the destination quickly; they must if they are not to fall into the forbidden comprehensive set. Because of this, as Ellison (2000) points out, they are not very reliable: routes that linger in a recurrent communicating class may be far more likely than direct routes to reach their destination. We will study such quasi-direct routes in Section 5.

Formally, define a *forbidden set* $W \subseteq Z$ for a path $a = (z_0, z_1, \dots, z_t)$ to be a set that the path does not touch except possibly at the beginning and end, that is, $z_1, \dots, z_{t-1} \notin W$.

DEFINITION 2 (Direct routes). Given an initial point $x \in Z$ and sets $B \subseteq W$, we call a nontrivial path a from x to B with forbidden set W a *direct route* if W is comprehensive and the path has finite resistance $r(a) < \infty$ (equivalently, positive probability for $\epsilon > 0$).

For each x, B and comprehensive W , there is a set of direct routes from x to B with forbidden set W , which we denote by A_{xBW} .⁸

We are interested in the following questions: How likely is the set of direct routes A_{xBW} , which paths in A_{xBW} are most likely, what are these paths like, and how long are they?

Results on direct routes

The intuition behind the results we present next is simple. Direct routes must hit the target without falling into a comprehensive set. This is hard, hence these routes have to be quick, and the quickest way is to make least resistance steps. This will be made precise in the following discussion.

First, to avoid triviality, we assume that $A_{xBW} \neq \emptyset$. An important observation is that there are typically many direct routes. Specifically, if there is a path $(z_0, z_1, \dots, z_t) \in A_{xBW}$ that contains a *loop*, that is, $z_\tau = z_{\tau'} \notin W$ for $\tau \neq \tau'$, then A_{xBW} is countably infinite since the loop can be repeated an arbitrary number of times. Notice also that if this loop has zero resistance, then the length of paths in A_{xBW} is without bound. Nevertheless, we shall see the expected length of such paths is quite short. For nonempty $A \subseteq A_{xBW}$, the first important fact proven in Appendix A is that $r(A) = \min_{a \in A} r(a)$ is well defined (and finite); it is the least resistance of any path in the set A .

The main result on direct paths characterizes their probability and length. The proof along with detailed bounds is given in Appendix A.

THEOREM 1. *If $A \subseteq A_{xBW}$ is nonempty, then $P_\epsilon(A|x) \sim \epsilon^{r(A)}$ and $E_\epsilon[t(a)|x, A] \sim 1$.*

In particular, positive resistance direct routes are not very likely to occur as ϵ gets small, yet they are unlikely to be terribly long in the sense that the expected length is

⁸The assumption that $B \subseteq W$ is without loss of generality. We can always define a forbidden set $W' = W \cup B$ without changing the set of direct routes. Note that for given x, B, W , the set of direct routes may be empty: it may be impossible to get to B without first hitting $W \setminus B$.

bounded independent of ϵ . Despite the fact that they are unlikely, these paths are important because they are needed to get from one recurrent communicating class to another. Intuitively, the reason these paths are short is that at each point along a direct route there is a zero-resistance path that leads to the forbidden set W . The more time that is spent along the route, the greater is the risk of falling into the forbidden set and failing to reach its destination. By contrast, we will see subsequently that the expected time spent in a recurrent communicating class goes to infinity as $\epsilon \rightarrow 0$.

The order of $P_\epsilon(A|x)$ established in [Theorem 1](#) directly implies the other facts characterizing direct routes.

COROLLARY 1. *Let $A = \{a \mid r(a) = r(A_{xBW})\}$ denote the least resistance paths in $A_{xBW} \neq \emptyset$. Then $\lim_{\epsilon \rightarrow 0} (P_\epsilon(A|x)/P_\epsilon(A_{xBW} \setminus A|x)) = \infty$.*

In other words, least resistance direct paths are far more likely than other direct paths. It is also the case that all least resistance direct paths have a probability similar to each other.

COROLLARY 2. *Let $A = \{a \mid r(a) = r(A_{xBW})\}$ and $a \in A$. Then $P_\epsilon(a|x)/P_\epsilon(A|x) \sim 1$.*

This completes the discussion of the basic results on direct routes.

5. TRANSITIONS BETWEEN RECURRENT COMMUNICATING CLASSES

We next study how the system is most likely to leave a recurrent communicating class and to which other class it is most likely to transit; then also how long it takes and where these paths spend their time along the way.

5.1 Quasi-direct routes

[Ellison \(2000\)](#) observes that being able to pass through every point in a recurrent communicating class may have a profound impact on the nature of the paths. The results of this section make this precise by characterizing quasi-direct routes that spend most of the time moving about without resistance within an initial class Ω_x .

We start again with an initial point $x \in Z$ in the recurrent communicating class Ω_x , a forbidden set $W \subseteq Z$, and a target set $B \subseteq W$. Notice that the definition in [Section 3](#) implies that direct routes from x to B with forbidden set W are not allowed to pass through all states in Ω_x , since the forbidden set W was assumed to be comprehensive. We now wish to relax that restriction and consider routes that are allowed to linger freely inside Ω_x .⁹ So we exclude Ω_x from the forbidden set, that is, we assume $W \cap \Omega_x = \emptyset$. Thus W cannot be comprehensive. However, we assume that W contains at least one point from every recurrent communicating class except for Ω_x . We then call W *quasi-comprehensive*.

⁹Notice that for the case of singleton Ω_x , this means the path may remain at x for some time or leave and return a number of times before hitting the target.

DEFINITION 3 (Quasi-direct routes). A nontrivial path a from $x \in Z$ to $B \subseteq W$ is a *quasi-direct route* if W is quasi-comprehensive and the path has finite resistance.

We denote the set of such paths by Q_{xBW} . As in the direct case, we assume the set Q_{xBW} is nonempty. Again we are interested in the structure of the paths in Q_{xBW} , in particular, which paths in Q_{xBW} are most likely, what do these paths look like, and how long are they?

Consider a path a in Q_{xBW} . The path originates at x and eventually hits B for the first time without first hitting W . The path may return to the start point x a number of times before finally departing and reaching the destination B . Consequently, it can be decomposed into a series of loops starting from x and returning to x , followed by a final departure or *exit path* to B . The loops start at x and return to x without hitting x or W in between; hence they are routes from x to x with forbidden set $W \cup \{x\}$. Since W is quasi-comprehensive, $W \cup \{x\}$ is comprehensive and we abbreviate these direct routes as $A_{xxW} \equiv A_{xx(W \cup \{x\})}$; this should not lead to confusion since W is quasi-comprehensive and so A_{xxW} has no other meaning. Following the loops, the exit path is a route from x to B that does not hit either x or W , that is, the forbidden set is again $W \cup \{x\}$ and so these are again direct routes that we abbreviate as $A_{xBW} \equiv A_{x(B \cup \{x\})}$. For $a \in Q_{xBW}$, we write $n(a)$ for the number of loops in a (it may be that $n(a) = 0$). Then if $n(a) > 0$, we can uniquely decompose $a \in Q_{xBW}$ as $a_1, a_2, \dots, a_{n(a)}, a^+$, where the $a_i \in A_{xxW}$ are the loops and $a^+ \in A_{xBW}$ is the exit path; if $n(a) = 0$ then $a = a^+$.¹⁰

Next we want a measure of the resistance of a quasi-direct path $a \in Q_{xBW}$. As we shall see, for such a path it is not the total resistance $r(a)$ that matters. What matters is the *peak resistance* $\rho(a) = \max\{r(a_i) \mid_{i=1}^{n(a)}, r(a^+)\}$, the greatest resistance of any of the loops or the exit path.

DEFINITION 4. The *least peak resistance* of a set $Q \subseteq Q_{xBW}$ is $\rho(Q) = \min_{a \in Q} \rho(a)$.

The next result shows that quasi-direct routes with least peak resistance consist of a least resistance exit path preceded by loops of weakly lower resistance.

THEOREM 2. If $a \in Q_{xBW}$ has least peak resistance so that $\rho(a) = \rho(Q_{xBW})$, then it has peak resistance equal to least exit resistance: $\rho(a) = r(a^+) = r(A_{xBW})$.

PROOF. Suppose $\rho(a) = \rho(Q_{xBW})$. From the definition $\rho(a) \geq r(A_{xBW})$, so the lemma can fail only if there is a path $\tilde{a} \in A_{xBW}$ for which $r(\tilde{a}) < \rho(a)$. But $\tilde{a} \in Q_{xBW}$, so this contradicts a having least peak resistance. $\square$

¹⁰Note that both A_{xxW} and A_{xBW} can contain loops from a point $y \neq x$ in Ω_x back to y provided we do not touch x in between. One can imagine alternative decompositions without this property, but this decomposition is just a tool for understanding how we get from x to B and technically our decomposition works well.

5.2 Leaving a recurrent communicating class

The following result (proved in [Appendix B](#)) plays a role in the theory of quasi-direct routes similar to that played by least resistance in the theory of direct routes in [Corollary 1](#).

THEOREM 3. *Let $A = \{a \mid \rho(a) = \rho(Q_{xBW})\}$ denote the least peak resistance paths in $Q_{xBW} \neq \emptyset$. Then $\lim_{\epsilon \rightarrow 0} (P_\epsilon(A|x)/P_\epsilon(Q_{xBW} \setminus A|x)) = \infty$.*

[Theorem 3](#) not only tells us the most likely routes from Ω_x to $\Omega \setminus \Omega_x$, by implication it also tells us where we are likely to leave Ω_x from and where we are likely to end up. First, since all states in Ω_x can be reached from x with no resistance, the path must leave Ω_x through a point $z \in \Omega_x$ from which the path to B is of least resistance among the direct routes from Ω_x to B with forbidden set $W \cup \Omega_x$, since leaving Ω_x through any other point would incur higher resistance. We call such z an *express exit*.

Next we consider where we end up. In case $B = W = \Omega \setminus \Omega_x$, which recurrent communicating class in $\Omega \setminus \Omega_x$ are we likely to move to? Let Ω_{-x}^{LP} be the collection of Ω_y in $\Omega \setminus \Omega_x$ for which there is a quasi-direct route from x to y of least peak resistance and let Ω_{-x}^{GP} be the remainder of $\Omega \setminus \Omega_x$. Let $P_\epsilon(\Omega_{-x}^j|x)$ denote the probability that starting at x , the first arrival at $\Omega \setminus \Omega_x$ is in Ω_{-x}^j for $j = \text{LP}, \text{GP}$. Then we have the following immediate corollary of [Theorem 3](#).

COROLLARY 3. *We have $\lim_{\epsilon \rightarrow 0} (P_\epsilon(\Omega_{-x}^{\text{LP}}|x)/P_\epsilon(\Omega_{-x}^{\text{GP}}|x)) = \infty$.*

To better understand what these results say about the actual dynamics of the system, recall that [Ellison \(2000\)](#) defines the *basin* of Ω_x as the set of states in Z for which there is a zero-resistance path to Ω_x and no zero-resistance paths to $\Omega \setminus \Omega_x$. Said otherwise, it is the set of states for which there is probability 1 in P_0 of reaching Ω_x . He also defines the *radius* as the least resistance of paths from Ω_x out of the basin. Focus on the case $B = W = \Omega \setminus \Omega_x$. [Theorem 2](#), together with the fact that there are zero-resistance paths from x to any other point $z \in \Omega_x$ and from outside the basin to $\Omega \setminus \Omega_x$ shows that the least peak resistance $\rho(Q_{xBW})$ is the same as the radius. [Theorem 3](#) shows that what matters for leaving the basin are the least peak resistance paths in the sense that paths from Ω_x to $\Omega \setminus \Omega_x$, which have peak resistance higher than the radius, are very unlikely. That includes both paths with a higher exit resistance than the radius and paths that have loops with higher resistance than the radius. [Corollary 3](#) shows that the recurrent communicating classes Ω_y that will be reached from Ω_x are very likely to be those for which the peak resistance is the same as the radius, which is to say that when we leave the basin on a direct route with resistance equal to the radius, there is then a zero-resistance path to Ω_y .

5.3 Expected length and visits of quasi-direct routes

We know from [Theorem 1](#) that transition paths in the direct route case are short. For paths that are allowed to remain in Ω_x we have the opposite result: these paths are quite

long. At this point it is convenient to focus on the case where B is reached with probability 1, that is, $P_\epsilon(Q_{xBW}|x) = 1$. We assure this by assuming that $B = W$.¹¹ Our goal is to show that Q_{xBW} has paths of expected length $\epsilon^{-r(A_{xBW})}$, that the fraction of time spent in Ω_x goes to 1, and that the absolute time spent outside of Ω_x goes to infinity. We will also show which kind of loops are likely to recur many times. The proofs of the results in this section may be found in [Appendix B](#).

The first result concerns the amount of time it takes to leave Ω_x and how much of that time is spent in Ω_x .

THEOREM 4. *Suppose that $B = W$. For $a = (a_1, a_2, \dots, a_{n(a)}, a^+) \in Q_{xBW}$, we let $a^- = (a_1, a_2, \dots, a_{n(a)})$ denote its loops and define $t^-(a)$ to be the amount of time along a^- spent outside of Ω_x .¹² Then*

$$E_\epsilon[t(a)|x, Q_{xBW}] \sim E_\epsilon[t(a^-)|x, Q_{xBW}] \sim \epsilon^{-r(A_{xBW})}$$

and

$$\lim_{\epsilon \rightarrow 0} E_\epsilon \left[\frac{t^-(a)}{t(a)} \middle| x, Q_{xBW} \right] = 0.$$

This says that quasi-direct paths including or excluding the exit path are long and spend most of their time in $\Omega(x)$.

The second result characterizes more exactly what happens while the system spends time outside of Ω_x during a quasi-direct route.

THEOREM 5. *Suppose that $B = W$. For $A \subseteq A_{xxW}$, let $M(a, A)$ be the number of loops of a that lie in A . For $A \subseteq A_{xxW}$*

$$E_\epsilon[M(a, A)|x, Q_{xBW}] \sim \epsilon^{r(A) - r(A_{xBW})}$$

if in addition $r(A) < r(A_{xBW})$ and $k \geq 0$

$$\lim_{\epsilon \rightarrow 0} P_\epsilon[M(a, A) > k|x, Q_{xBW}] = 1.$$

Also let $A_{xxW}[t]$ be the set of loops that spend at least t consecutive periods outside of Ω_x . If there is a path $a^0 \in A_{xxW}$ that contains a zero-resistance loop not touching Ω_x with $0 < r(a^0) < r(A_{xBW})$, then for any $k > 0$,

$$\lim_{\epsilon \rightarrow 0} P_\epsilon(M(a, A_{xxW}(kt(a^+))) > k|x, Q_{xBW}) = 1.$$

¹¹Recall that the paths in Q_{xBW} by definition have positive probability of reaching B without touching W along the way; when $B = W$, there is no other way to reach B , so this probability becomes 1.

¹²This first result is an extension of [Ellison's \(2000\)](#) result that the waiting time for leaving Ω_x is of order ϵ^{-r} , where r is the radius. Ellison mentions both an upper and lower bound in the text, but we have been unable to locate his proof of the lower bound. The extension here is that W can be a general quasi-comprehensive set, for example, it might include states in the basin of Ω_x .

The first two statements say that if there is some $a^0 \in A_{xxW}$ with $0 < r(a^0) < r(A_{xBW})$, then this loop will occur many times even though the fraction of the time spent outside Ω_x in such loops must be small. Alternatively, we can say it this way: as $\epsilon \rightarrow 0$, loops in A_{xxW} that have resistance strictly less than the radius occur an arbitrarily large number of times before we leave the basin of x and those that have a resistance strictly greater than the radius have a vanishing small probability of occurring before we exit the basin. The third result says that it will often be the case that a will spend more than k times as long outside of $\Omega(x)$ as it takes to get to the final destination following the exit path.

On a more technical note, the first result describes the expected number of occurrences, while the second result describes the realized number of occurrences of loops in A . The difference between the two is that the amount of time before leaving the basin is random. It could be that with very high probability, the number of occurrences is small (less than k say), and in those rare cases where the length of time before leaving the basin is very large, the number of occurrences is very large. In this case the expected number of occurrences may grow as ϵ gets smaller, while the probability of seeing more than k occurrences remains unchanged or even falls. The second part of [Theorem 5](#) shows that this cannot happen.

6. THE BIG PICTURE

Reconsider the dynamics of P_ϵ . Starting at any point x , by [Theorem 1](#), we move quickly to one of the recurrent communicating sets Ω_y . Once there, by [Theorem 4](#), it is a long time before we reach a different Ω_z , and most of that time is spent in Ω_y . One question we now address is what the dynamics look like during the long period when we are in Ω_y . When we do finally leave Ω_y , by [Theorem 2](#) we move quickly to the next Ω_z , and it is most likely the recurrent communicating set that has least exit resistance from Ω_y . The second question we will address is, over the longer run, how much time do we spend in the different recurrent communicating sets in Ω ? To this end, we assume in this section that P_ϵ is ergodic for $\epsilon > 0$ and denote by μ_ϵ the unique ergodic distribution of the process.

The big picture that we will break down over the next subsections can be visualized by thinking of astronomy. At the bottom level are planets, which correspond to recurrent communicating classes. Movements within recurrent communicating classes can be thought of as moving around on the planetary surface, that is, relatively quick. Recurrent communicating classes are surrounded by states in their “inner basin” that are tightly bound to them like moons around a planet. There are also a few states that are either far from recurrent communicating classes or bound to several of them. For these only we cannot give precise bounds on the ergodic probabilities; we may think of them as comets. Recurrent communicating classes in turn are grouped into “circuits”—think of those as solar systems: movement within a solar system being much more rapid than movement between solar systems. These circuits—solar systems—are grouped into higher order circuits—galaxies—and the “galaxies” in turn are grouped to even higher level circuits and so forth. We give tight bounds for the relative ergodic probabilities for elements within a circuit: planets within a solar system, solar systems within a galaxy, and so forth. In the end, all these circuits are contained in a single universe, and this will give us the relative ergodic probabilities of all recurrent communicating classes.

We should emphasize that while the method of circuits can be turned loose on arbitrary models to determine the structure and relative probability of recurrent communicating classes, it can also be useful in building models. Some structures are easier to analyze than others. For example, in the hegemony model of [Levine and Modica \(2013a\)](#) all recurrent communicating classes are in the same circuit, so it is trivial to analyze their relative resistances: it is given by the differences of least resistances to leaving, or more simply in that context, hegemonies can be ranked by their state power in their worst state, and hegemonies with higher state power are relatively much more likely in the ergodic distribution than those with lesser state power. At the next level, if we can establish assumptions that bind recurrent communicating classes into groups where all the groups lie in a single circuit, then within each group relative resistances are determined by exit resistances, and the relative resistances between groups is determined by the exit resistances from group to group. More broadly, our intuition may tell us what the structure of circuits “ought” to look like so that it would be natural to focus on assumptions that lead to that type of structure.

6.1 Inside recurrent communicating classes: On the surface of the planet

When $\epsilon = 0$ we cannot move between recurrent communicating classes, but we have well defined and ergodic dynamics within each class. Moreover, these dynamics are fast in the sense that they are independent of ϵ . Hence if we are interested in the approximate probability of events within a class, we should consider paths of bounded length. For such events, the probabilities in P_0 are much the same as for P_ϵ for ϵ small. Specifically, we state the following theorem.

THEOREM 6. *Let A_1 and A_2 be any collections of paths of bounded length starting at $x \in \Omega_x$ and for which $P_0(A_2|x) > 0$. Then*

$$\lim_{\epsilon \rightarrow 0} \frac{P_\epsilon(A_1|x)}{P_\epsilon(A_2|x)} = \frac{P_0(A_1|x)}{P_0(A_2|x)}.$$

PROOF. Since the probabilities are defined by finite sums of finite products of the transition probabilities $P_\epsilon(z|y)$, and the length of the sums and the products are bounded independent of ϵ , the result follows immediately from the assumption that $\lim_{\epsilon \rightarrow 0} P_\epsilon(z|y) = P_0(z|y)$. $\square$

In particular, since paths that lie entirely within Ω_x have probability 1 in P_0 given x , the probability of sets of paths of finite length within Ω_x is roughly the same in P_ϵ as in P_0 when ϵ is small.

The other important characteristic of Ω_x is the amount of time spent at different states. Notice that if we restrict the state space to Ω_x , then P_0 is an ergodic Markov process on that space, and so has a unique and strictly positive ergodic distribution $\bar{\mu}_0(y)$, where $\sum_{y \in \Omega_x} \bar{\mu}_0(y) = 1$. Notice in particular that if $y \in \Omega_x$, the ratio $\bar{\mu}_0(x)/\bar{\mu}_0(y)$ is well defined and finite. We can relate this to the ratio of stationary probabilities $\mu_\epsilon(x)/\mu_\epsilon(y)$ for the process when $\epsilon > 0$. In [Appendix C](#) we show that the following theorem holds.

THEOREM 7. If $y \in \Omega_x$, then

$$\lim_{\epsilon \rightarrow 0} \frac{\mu_\epsilon(x)}{\mu_\epsilon(y)} = \frac{\bar{\mu}_0(x)}{\bar{\mu}_0(y)}.$$

6.2 Basins and ranges: Moons and comets

Recall that the basin of Ω_x consists of the states for which there is probability 1 in P_0 of reaching Ω_x . For the purpose of relating states to recurrent communicating classes, it is useful to define three variations on the basin. First observe that from any point outside the basin there is positive P_0 probability of reaching $\Omega \setminus \Omega_x$, that is, from outside the basin there is a zero-resistance path to $\Omega \setminus \Omega_x$. Therefore, the radius—the least resistance of leaving the basin—is also the least resistance to get to $\Omega \setminus \Omega_x$.

DEFINITION 5. The *outer range* of Ω_x is the set of states for which there is a zero-resistance path to Ω_x . The *inner range* of Ω_x is the set of states y in the outer range that can also be reached from x with resistance not larger than the radius of Ω_x .¹³ If we further require that the resistance of being reached is strictly less than the radius, we have the *inner basin*.

The outer range is the largest set of states affiliated with Ω_x in the sense that any states outside the outer range will never get to Ω_x when $\epsilon = 0$. The outer range contains the basin, but unlike the basin, the outer range does not require reaching Ω_x with probability 1 when $\epsilon = 0$, and a point can be in the outer range of different recurrent communicating classes.¹⁴ By contrast, the basins of different recurrent communicating classes must be disjoint.

The inner basin is a subset of the inner range by definition. It is also a subset of the basin, since if a point in the inner basin were not in the basin, we could go from x to $\Omega \setminus \Omega_x$ with resistance strictly less than the radius, which is impossible by definition. The states in the inner basin are the states most tightly affiliated with Ω_x : this is shown by Theorem 4, which says that these states will be hit many times before moving on to the next recurrent communicating class.

There is no firm relationship between the basin and the inner range. The basin may contain states further from Ω_x than the radius, so states not in the inner range. The inner range contains states at a distance equal to the radius, and some of these states must have zero-resistance paths to other recurrent communicating classes so that they cannot be part of the basin. Nevertheless, states in the inner range are still “close” to Ω_x : they are part of least peak resistance paths to other recurrent communicating classes, and Theorem 4 says the expected number of times they will be hit before moving on is positive. By contrast, the states that are in the basin but not the inner range are “far” from Ω_x and this can be seen precisely in Theorem 3, which shows that these states

¹³If this is true, then there is also a direct route with forbidden set $W = (\Omega \setminus \Omega_x) \cup \{x\} \cup \{y\}$ with this property. Basically, within the basin it does not matter whether least resistance is measured along direct routes or all routes since it is not possible to pass through $\Omega \setminus \Omega_x$ while remaining in the basin.

¹⁴In the introductory example, the state $\{2\}$ is in the outer range of both Ω_0 and Ω_5 .

are unlikely to be reached from Ω_x prior to reaching another recurrent communicating class.

States in the inner basin are like “moons” tightly bound to the recurrent communicating class, while states in the outer range but not the inner basin are more like “comets” that are not tightly bound to any recurrent communicating class.

Let $r^D(x, y)$ be the least resistance of any direct path from $x \in \Omega_x$ to y that does not pass through any recurrent communicating class other than Ω_x , that is, we define $r^D(x, y) \equiv r(A_{xyW})$ with $W = (\Omega \setminus \Omega_x) \cup \{x\} \cup \{y\}$. Take $W = \Omega \setminus \Omega_x$ and define the radius¹⁵ of Ω_x to be $r^0(\Omega_x) = r(Q_{xWW})$. The following theorem is shown in [Appendix C](#).

THEOREM 8. *If $x \in \Omega_x$ and $r^D(y, x) = 0$, then*

$$\frac{\mu_\epsilon(y)}{\mu_\epsilon(x)} \sim \epsilon^r,$$

where $\min\{r^D(x, y), r^0(\Omega_x)\} \leq r \leq r^D(x, y)$.

For the comets—states that are not in the inner range of any recurrent communicating class, that is, states that are in the outer range of one or more recurrent communicating class, but are “hard” to reach—[Theorem 8](#) gives bounds for the ergodic probabilities: they cannot be too likely. For the moons—states in the inner range—the bound is tight, and says that their relative probability of occurring is inversely proportional to the least resistance of getting to them from Ω_x .

COROLLARY 4. *If y is in the inner range of $x \in \Omega_x$, then*

$$\frac{\mu_\epsilon(y)}{\mu_\epsilon(x)} \sim \epsilon^{r^D(x, y)}.$$

6.3 Recurrent communicating classes: Solar systems, galaxies and beyond

Consider again the overview of the dynamics of P_ϵ : starting at any point x , we move quickly to one of the recurrent communicating sets Ω_y ; once there it is a long time before we reach a different Ω_z and most of that time is spent in Ω_y . When we do finally leave Ω_y , we move quickly—and directly, in our sense—to the next Ω_z and it is most likely the recurrent communicating set that has least exit resistance from Ω_y . Proceeding in this way, we get a sequence of recurrent communicating sets Ω_i connected by least exit resistances.¹⁶ Since the set Ω of recurrent communicating classes in P_0 is finite, eventually this sequence must have a cycle.

More general than the notion of a cycle, we introduce the notion of a circuit. The defining property of a circuit is that between any two of its states there is a path within the circuit with least resistance transitions. We start by defining circuits in Ω . For $\Omega_x, \Omega_y \in \Omega$, define transition resistance $r^0(\Omega_x, \Omega_y) = \min\{r^D(x, y) \mid y \in \Omega_y\}$ that is the least resistance of any direct path from x to Ω_y not touching $\Omega \setminus \Omega_x$; least resistance out of Ω_x is then defined as $r^0(\Omega_x) = \min_{\Omega_y \in \Omega \setminus \Omega_x} r^0(\Omega_x, \Omega_y)$ —[Ellison’s \(2000\)](#) radius of Ω_x .

¹⁵As noted above, this is different from [Ellison’s \(2000\)](#) definition but is equivalent.

¹⁶We could equally well say radius.

DEFINITION 6 (Circuits). A set $\Omega_x^1 \subseteq \Omega$ is a *circuit* if for any pair $\Omega_1, \Omega_y \in \Omega_x^1$, there is a path $(\Omega_1, \Omega_2, \dots, \Omega_n)$ in Ω_x^1 with $\Omega_n = \Omega_y$ and $r^0(\Omega_{\tau-1}, \Omega_\tau) = r^0(\Omega_{\tau-1})$ for $\tau = 2, 3, \dots, n$.

Our basic observation is that once we reach a circuit, we remain within the circuit for a long time before going to another circuit. We first compare states within a circuit. Since the probability of leaving Ω_x is of order $\epsilon^{r^0(\Omega_x)}$, the expected length of any visit to Ω_x is $1/\epsilon^{r^0(\Omega_x)}$. We might then expect that the amount of time we spend at Ω_x is roughly $\epsilon^{r^0(\Omega_y) - r^0(\Omega_x)}$ as long as the amount of time we spend at Ω_y . In Appendix S.D, available in a supplementary file on the journal website, <http://econtheory.org/supp/1978/supplement.pdf>, we show that this is indeed true.

THEOREM 9. *If the recurrent communicating classes Ω_x and Ω_y are in the same circuit, then*

$$\frac{\mu_\epsilon(x)}{\mu_\epsilon(y)} \sim \epsilon^{r^0(\Omega_y) - r^0(\Omega_x)}.$$

We next ask how long do we actually spend in a circuit over Ω ? We stay in Ω_x roughly $\epsilon^{-r^0(\Omega_x)}$ periods; and the probability of going to a fixed $\Omega_y \neq \Omega_x$ out of the circuit is of order $\epsilon^{r^0(\Omega_x, \Omega_y)}$. Hence the probability of going to Ω_y during a visit to Ω_x is of order $(1/\epsilon^{r^0(\Omega_x)})\epsilon^{r^0(\Omega_x, \Omega_y)}$. For this to occur with very high probability, the number of visits to Ω_x must be roughly k_x , where $k_x(1/\epsilon^{r^0(\Omega_x)})\epsilon^{r^0(\Omega_x, \Omega_y)} = 1$; that is, $k_x = 1/\epsilon^{r^0(\Omega_x, \Omega_y) - r^0(\Omega_x)}$. Following Ellison (2000), we define the *modified resistance from Ω_x to Ω_y* as $R^0(\Omega_x, \Omega_y) = r^0(\Omega_x, \Omega_y) - r^0(\Omega_x)$. Then the number of visits is least for the element Ω_z in the circuit that has minimum $R^0(\Omega_z, \Omega_y)$ over $\Omega_y \notin \Omega_x^1$. This is the most likely (actually least modified resistant) exit from the circuit. Also, it will exit to a circuit that is easiest to reach. This in turn suggests that we can form circuits of circuits using modified resistances as the measure of resistance in going from one circuit to another. The system moves between circuits of circuits in a longer time horizon. Moreover, as we have seen, the crossings between circuits are direct routes; hence we will define resistance in terms of such paths. We spell out these ideas next.

The procedure we will describe is essentially the same as that employed by Cui and Zhai (2010) to compute the stochastically stable state.¹⁷ Here we employ the procedure to determine the relative probabilities of recurrent communicating classes.

We build circuits recursively starting from $\Omega^0 \equiv \Omega$. Assuming Ω has $N_\Omega \geq 2$ elements, we observe in Appendix S.D that there is at least one circuit that is nontrivial in the sense of having at least two elements, and that every singleton element is trivially a circuit. Hence we can form a nontrivial partition of Ω^0 into circuits, and we denote by Ω^1 the collection of elements of this partition—so an element of Ω^1 is a circuit on $\Omega = \Omega^0$. Given two distinct elements $\Omega_x^1, \Omega_y^1 \in \Omega^1$, define transition resistance $r^1(\Omega_x^1, \Omega_y^1) =$

¹⁷There are a number of known algorithms for computing the stochastically stable state related to that of Cui and Zhai (2010). Hasker (2014) has a good review of the literature. Cui and Zhai (2010) in fact use the terminology “cycle,” but as that seems misleading, we prefer the term “circuit,” since movement within a circuit will not necessarily be a cycle.

$\min\{R^0(\Omega_x, \Omega_y) \mid \Omega_x \in \Omega_x^1, \Omega_y \in \Omega_y^1\}$, least outgoing resistance $r^1(\Omega_x^1) = \min\{r^1(\Omega_x^1, \Omega_y^1) \mid \Omega_y^1 \in \Omega^1 \setminus \Omega_x^1\}$, and modified resistance $R^1(\Omega_x^1, \Omega_y^1) = r^1(\Omega_x^1, \Omega_y^1) - r^1(\Omega_x^1)$. A set in Ω^1 is a circuit if, as before, between any two of its states there is a path where each transition has least outgoing resistance (according to r^1 of course). Continuing, as long as Ω^{k-1} has more than one element we partition it into circuits, call Ω^k the resulting collection of elements, and for $\Omega_x^k, \Omega_y^k \in \Omega^k$ define $r^k(\Omega_x^k, \Omega_y^k) = \min\{R^{k-1}(\Omega_x^{k-1}, \Omega_y^{k-1}) \mid \Omega_x^{k-1} \in \Omega_x^k, \Omega_y^{k-1} \in \Omega_y^k\}$, $r^k(\Omega_x^k) = \min\{r^k(\Omega_x^k, \Omega_y^k) \mid \Omega_y^k \in \Omega^k \setminus \Omega_x^k\}$, and modified resistance $R^k(\Omega_x^k, \Omega_y^k) = r^k(\Omega_x^k, \Omega_y^k) - r^k(\Omega_x^k)$. Note that since each partition is nontrivial, this construction has at most N_Ω layers before the partition has a single element and the construction stops: $k \leq N_\Omega$.

The crucial function at each layer k turns out to be the *modified radius* of $x \in \Omega_x$ of order k , defined by

$$\bar{R}^k(x) = \sum_{\kappa=0}^k r^\kappa(\Omega_x^\kappa),$$

where $\Omega_x^0 = \Omega_x$ and for each $\kappa > 0$, the element $\Omega_x^\kappa \ni \Omega_x^{\kappa-1}$. It is a measure of the difficulty of traveling far from x and has an intuition similar to that of Ellison's (2000) modified co-radius. We show in Appendix S.D that the following theorem holds.

THEOREM 10. *Let k be such that $\Omega_x^k = \Omega_y^k$, that is, x and y are in the same circuit in Ω^k . Then*

$$\frac{\mu_\epsilon(x)}{\mu_\epsilon(y)} \sim \epsilon^{\bar{R}^{k-1}(y) - \bar{R}^{k-1}(x)}.$$

It is useful to define the difference between the modified radii $\bar{R}^{k-1}(y) - \bar{R}^{k-1}(x)$ as the *relative ergodic resistance* of x over y . The theorem says that the relative probabilities within a circuit are proportional to ϵ to the power of the relative ergodic resistance. Note here the implication: if x and y are in the same circuit in Ω^k , then of course they have to be in the same circuit in any Ω^κ for $\kappa > k$. Hence in this case $\bar{R}^{\kappa-1}(y) - \bar{R}^{\kappa-1}(x) = \bar{R}^{k-1}(y) - \bar{R}^{k-1}(x)$.

Traditionally, interest has focused on the states x that have ergodic probabilities that are bounded away from zero—the *stochastically stable* states. We can see from [Theorem 10](#) that in any circuit these states must have nonpositive relative ergodic resistance over any other state and strictly negative over any state that is not also stochastically stable. Since this must be true in any circuit, it must be true in the top level circuit that contains all recurrent communicating classes, that is to say, the Ω^k where k is the highest level of the filtration where the partition is a singleton. We thus have the following corollary.

COROLLARY 5. *The stochastically stable states are exactly those with the highest values of $\bar{R}^{k-1}(x)$, where k is the level at which the partition into circuits is a singleton.*

6.4 An example

To illustrate the application of [Theorem 10](#), let us give a complete analysis of the case where Ω has three elements.¹⁸ Note that there are nine trees on three states, so the analysis by means of trees is already difficult. For simplicity let us make the generic assumption that no two resistances or sums or differences of resistances are equal.

There are two cases: either there is a single circuit or there is one circuit consisting of two states and an isolated point. The case of a single circuit is trivial: in this case, in [Theorem 10](#), $k = 1$, and since $\bar{R}^0(x) = r^0(\Omega_x)$, the relative ergodic resistances are given simply by the differences in least outgoing resistances between the three states, and the stochastically stable state is the point with least outgoing resistance.

Take then, the case of Ω with one two-point circuit and an isolated point. Denote by Ω_x and Ω_y the two states on the circuit and denote by Ω_z the remaining point. Then $k = 2$ in the theorem. Assume without loss of generality that $r^0(\Omega_x) > r^0(\Omega_y)$ so that, within the circuit, Ω_x is relatively more likely. Notice that $r^0(\Omega_x, \Omega_y) < r^0(\Omega_x, \Omega_z)$ and $r^0(\Omega_y, \Omega_x) < r^0(\Omega_y, \Omega_z)$ since Ω_x and Ω_y are on the same circuit; this also implies $r^0(\Omega_x) = r^0(\Omega_x, \Omega_y)$ and $r^0(\Omega_y) = r^0(\Omega_y, \Omega_x)$. Turning to the recursion, we need to work out the least modified resistances. Let $\Omega_x^1 = \{\Omega_x, \Omega_y\}$ be the circuit and let $\Omega_z^1 = \Omega_z$ be the isolated point. Then $r^1(\Omega_z^1, \Omega_x^1) = \min\{r^0(\Omega_z, \Omega_x) - r^0(\Omega_z), r^0(\Omega_z, \Omega_y) - r^0(\Omega_z)\} = 0$, while $r^1(\Omega_x^1, \Omega_z^1) = \min\{r^0(\Omega_x, \Omega_z) - r^0(\Omega_x), r^0(\Omega_y, \Omega_z) - r^0(\Omega_y)\}$. Hence $\bar{R}^1(z) = r^0(\Omega_z)$, which is just the radius of Ω_z , while

$$\begin{aligned}\bar{R}^1(x) &= r^0(\Omega_x) + \min\{r^0(\Omega_x, \Omega_z) - r^0(\Omega_x), r^0(\Omega_y, \Omega_z) - r^0(\Omega_y)\} \\ &= \min\{r^0(\Omega_x, \Omega_z), r^0(\Omega_x) + r^0(\Omega_y, \Omega_z) - r^0(\Omega_y)\},\end{aligned}$$

which is to say exactly what [Ellison \(2000\)](#) defines as the modified co-radius of Ω_z .¹⁹ The relative ergodic resistance of x over y is therefore $r^0(\Omega_z) - \bar{R}^1(x)$, while the relative ergodic resistance of y can be recovered from the relative ergodic resistance of y over x , which is just $r^0(\Omega_x) - r^0(\Omega_y)$. With respect to stochastic stability, we see that Ω_z is stochastically stable if and only if its radius $r^0(\Omega_z)$ is greater than its co-radius $\bar{R}^1(x)$, which is [Ellison's \(2000\)](#) sufficient condition, and otherwise Ω_x is not stochastically stable. In short, the entire ergodic picture comes down to computing three numbers: the radius and co-radius of Ω_z and the difference between the radii of Ω_x and Ω_y .

7. THE FALL OF HEGEMONIES

We now discuss how our results can be used to interpret historical facts concerning sequences of long-run social events of small probability. This is the natural field of application of our theory, which concerns transitions along paths whose steps are each quite

¹⁸See also [Hasker \(2014\)](#).

¹⁹In fact, the modified co-radius is defined as the larger of $\bar{R}^1(x)$ and $\bar{R}^1(y) = \min\{r^0(\Omega_y, \Omega_z), r^0(\Omega_y, \Omega_x) + r^0(\Omega_x, \Omega_z) - r^0(\Omega_x, \Omega_y)\}$. However, $r^0(\Omega_x, \Omega_z) > r^0(\Omega_y, \Omega_x) + r^0(\Omega_x, \Omega_z) - r^0(\Omega_x, \Omega_y)$ and $r^0(\Omega_x, \Omega_y) + r^0(\Omega_y, \Omega_z) - r^0(\Omega_y, \Omega_x) > r^0(\Omega_y, \Omega_z)$ imply $\bar{R}^1(x) \geq \bar{R}^1(y)$.

unlikely to occur. We will focus in particular on the fall of the Qing dynasty in twentieth century China. We use a variation of the model of [Levine and Modica \(2013a\)](#) and [Levine and Modica \(2013b\)](#), identifying the fall of a hegemonic society with its progressive loss of land to a different society.²⁰ We emphasize that this application is limited in scope: the goal is to illustrate how results on transitions lead to interesting predictions. The sensitivity of the predictions to assumptions and the validity of those predictions across a broad range of data are of independent interest but are beyond the scope of this paper and are discussed in [Levine and Modica \(2013b\)](#).

There are a finite number J of societies. In each period t , each society j has one of a finite number of internal states $\xi_{jt} \in \Xi_j$. These states evolve according to a fixed Markov process $\Pi_j(\xi_{jt}|\xi_{j,t-1}) > 0$ independent of ϵ so that all transitions are possible. External forces such as disease, climate, other real shocks to productivity, the interference of outsiders who are protected themselves by geographical barriers, or superior technology can lead to changes in the internal state; the state may also represent changes in the internal structure of institutions. A good example of the Π_j process within a given society can be found in [Acemoglu and Robinson \(2001\)](#): there external shocks in the form of recessions drive changes in institutions whereby the voting franchise is extended or contracted. The ability of a society to resist and influence other societies is indexed by “state power” $\gamma_j(\xi_{jt})$. Societies may or may not satisfy incentive constraints: we represent this by a stability index $b_j \in \{0, 1\}$ with 1 indicating stability, where societies violating incentive constraints are thought to be unstable. We assume that the strongest unstable society with the most favorable value of ξ_j is stronger than the strongest stable society in its least favorable value of ξ_j , that is, $\max_{j|b_j=0, \xi_{jt} \in \Xi_j} \gamma_j(\xi_{jt}) > \max_{j|b_j=1} \min_{\xi_{jt} \in \Xi_j} \gamma_j(\xi_{jt})$. This reflects the idea that unstable societies face weaker incentive constraints.

Societies compete over a single resource called land. Each society j holds an integral number of units of land L_j , where there are L units of land in total. A state z is a list of landholding and real shocks of the different societies, $z = (L_1, \xi_1, L_2, \xi_2, \dots, L_J, \xi_J)$. Land changes ownership between societies due to conflict. We assume that at most 1 unit of land changes hands each period. The probability that society j loses a unit of land is given by a conflict resolution function with resistance $r_j(z) < \infty$ to j losing 1 unit of land; that is, the resistance of a transition from z to a state where j has one less unit of land. If j loses land, the probability the land goes to society k has *land gain resistance* $\lambda_{jk}(z)$, where $\lambda_{jj} = \infty$, but if $k \neq j$, then $\lambda_{jk} < \infty$.

Conceptually, in this model there are two distinct types of societies: *active societies* that have positive landholdings and *inactive societies* that do not. Inactive societies represent templates for societies that might exist but do not exist currently: an inactive society may become active because when an active society loses land, the land may be lost to an inactive society, that is, the loss of land may represent experimentation with new institutions.

²⁰[Levine and Modica \(2013a, 2013b\)](#) use a model in which the conflict is distinct from learning. The model here simplifies that model by having learning take place on a single unit of land at a time when a society is “unstable.” This greatly simplifies presentation of the model without any important consequences for the results. The model here otherwise weakens the assumptions in those papers.

We make several specific assumptions about the conflict resolution and land gain resistance. We assume that r and λ depend only on the landholdings, state power, and stability of the different societies. Since they are just templates for societies, we assume that the state power of inactive societies does not matter. We assume that conflict resolution resistance is monotone, so that the probability of j losing land decreases with its own state power and land, and increases with that of other societies. In particular, we assume that unstable societies always have zero resistance to losing a unit of land: if incentive constraints are not satisfied, individuals experiment with different actions and societies experiment with different institutions (albeit just on a single unit of land at a time). For stable societies, we assume the resistance is strictly monotone when nonzero and that the weakest society with positive landholding has zero resistance to losing a unit of land. Finally, we assume for given landholding that resistance is greater when facing more than one opponent who has positive landholding than when all enemy land is in the hands of the strongest landholding opponent—also that this is strict if resistance is positive. With respect to land gain resistance, we assume that if $k \neq j$ and $L_k > 0$ then $\lambda_{jk}(z) = 0$, that is, active societies all have zero resistance to gaining land.

7.1 *The rise, the fall, and warring states*

We apply the same outline of analysis to this model as we did to the simple example of [Section 2](#): we look first for the recurrent communicating classes, then for the relationship between them, then the basins, then the relationship between the basins and recurrent communicating classes, and finally examine the transitions.

1. Recurrent communicating classes for $\epsilon = 0$ A *hegemony* is a single society that controls all the land. The assumption that the weakest society with positive landholding has zero resistance to losing a unit of land plays a key role in determining the recurrent communicating classes. It implies that there is a zero-resistance path from every non-hegemonic state to a hegemonic state: monotonicity implies that losing land cannot increase resistance, so the weakest society keeps losing land until hegemony is established. Second, there is a zero-resistance path from any unstable hegemony to a stable hegemony: by assumption, the unstable hegemony has zero resistance to losing land and by symmetry, the first unit of land lost has zero resistance to being “taken” by a stable society, which once it becomes active continues to have zero resistance to taking over the next unit of land and so forth.

There are three cases, depending on what the stable hegemonies are like. It may be that the stable hegemony has so little state power that it also has zero resistance to losing land. In this case, there is a zero-resistance path from a hegemony to any state in which there are no more than two active societies (and, in particular, from any hegemony to any other), and from there to other states without a hegemonic society, so that there is only one recurrent communicating class and it is “large” in the sense that it includes within it unstable as well as stable societies and societies with all levels of state power. Second, it may be that there is a single stable hegemony that has the greatest state power and that only this hegemony is strong enough to resist losing land. In this case, that

single hegemony Ω_x is the only recurrent communicating class. If $y \in \Omega_x$, then $j(x)$ can denote the single stable society $j(x)$ that controls all the land, while the different states $y \in \Omega_x$ correspond to different shocks $\xi_{j(x)}$. Finally, it may be that two or more stable societies are strong enough to have resistance when hegemonic. These different societies may have the same or different state power. Regardless, these societies each constitute a recurrent communicating class. Our interest here is in the third case—about how hegemonies fall—that is, how we move from a hegemony Ω_x to $B = \Omega \setminus \Omega_x$, $W = B$.

2. Relation between the classes for $\epsilon > 0$ Again we ask, starting in one recurrent communicating class, which comes next, how long will it take, and relatively how much time is spent at each recurrent communicating class?

The first step is the standard one of understanding the radius and least resistance paths out of the basin that by [Corollary 3](#) we know are the most likely ways of leaving the basin. Because of monotonicity, a single invader with the greatest state power in the most favorable state always has the least resistance to gaining a unit of land from the hegemon. Hence a path in which such an invader repeatedly takes land from the hegemon is a least resistance direct route out of the basin. There is a threshold level of land such that once the hegemon loses this amount of land, it loses resistance: because this “optimal” invader is necessarily at least as strong as the hegemon and because the weaker of the two societies always has zero resistance to losing land, this threshold is no more than 50% of the land.

Next we observe that we have assumed that unstable societies can generate more state power than any stable society. This means that the optimal invader must be unstable. We refer to the unstable society that generates the greatest state power as *zealots*. These are the optimal invaders.

Notice that once the basin has been exited, zealots can continue to gain land with zero resistance until there is a hegemony of zealots. At this point, since zealots are unstable and have zero resistance to losing land themselves, their hegemony is transient, and any stable society can enter and take all the land away from the zealots until they themselves form a hegemony. Hence, from any $\Omega_x \in \Omega$, there is a least resistance path to any other $\Omega_y \in \Omega$. That means that all the recurrent communicating classes are in the same circuit and so by [Theorem 9](#) their relative probabilities simply depend on the differences in their radii. Since the radius is determined entirely by state power and is strictly increasing in state power, this means that hegemonies with greater state power are far more likely in the long run than societies with less state power. Notice how the (very reasonable) assumption that greater power is generated by ignoring incentive constraints than by satisfying them leads to this very simple long-run dynamic.

By the earlier work of [Ellison \(2000\)](#) and [Theorem 4](#), we know that the waiting time for a hegemony to fail is determined by the radius and, hence by the state power. Hegemonies with greater state power are more durable.

3. Basins The states in the basin of a hegemony Ω_x are simply those states in which the society has positive resistance to losing land for all internal states $\xi_{j(x)}$. Since this means that $j(x)$ is not the weakest society, it follows that there is zero resistance to some other society losing land, and since $j(x)$ has zero resistance to gaining that land and continues

to have positive resistance to losing land when it increases its landholding, this means a probability 1 path to the hegemony when $\epsilon = 0$. The inner basin is just that subset of the basin that can be hit with less resistance than the radius: this is typically a much smaller set than the basin.²¹

The outer range consists of states that are in the basin but have resistance strictly greater than the radius, plus a large collection of states that have zero resistance both to returning to the hegemony and to some other recurrent communicating class. This is true of any state in which an unstable society is the strongest (in its strongest internal state), such as those states reached after invasion by zealots.²² The inner range just adds those states in the outer range that have resistance equal to the radius.

4. Relation between basins and classes By [Theorem 4](#), during the time before reaching a new class most of the time will be spent in the current recurrent communicating class, but there will be a large number of periods during which the system will move to states in the inner basin and back. From [Theorem 8](#), the relative amount of time spent at these states compared to the time spent in the recurrent communicating class is an exponential function of the difference between the resistance of reaching the point and the radius; further states are less likely.

5. Transitions How we get from one recurrent communicating class to another is covered in [Theorem 3](#). This says that when ϵ is small, it is nearly certain that this transition will take place along a least peak resistance path, and we want to describe such paths in the model at hand: such a path may leave the recurrent communicating class and return any number of times, but during each departure from the recurrent communicating class the resistance encountered can be no more than the radius.

Notice that from an empirical point of view we probably do not know what the inner basin looks like, but we may have a pretty good idea of a bound on the amount of land that the hegemon needs to be in the inner basin. For example, we may think that if the hegemon loses 30% or less of its land it remains “safe” in the sense that it still has resistance. This shows some of the strength of using arbitrary quasi-comprehensive sets W . We can simply take the forbidden set and the target set $B = W$ to be the set of states in which the hegemon controls at least 70% of the land: all the theorems about least peak resistance apply equally well to this W , except that instead of radius, we simply say “resistance to losing 30% of land,” which has more empirical content.

²¹A point is in the basin but not the inner basin if the hegemony is reached with probability 1 when $\epsilon = 0$ but the resistance to reaching the point is greater than the radius. So, for example, if a sufficiently weak opponent occupies a large fraction of the land, then it will be overcome by the hegemony when $\epsilon = 0$ with probability 1, yet starting at the hegemony, the resistance to such a weak opponent grabbing so much land can be much greater than the resistance to determined zealots taking enough land to put an end to the hegemony.

²²As we remarked above, because the unstable society is strongest, there is zero resistance to the unstable society taking over. However, once the unstable society has taken over, it has zero resistance to losing all its land to any hegemony. In other words, from such a state there is a zero-resistance path to every recurrent communicating class. Consequently these states are in the outer range, yet they are certainly not in the basin.

Prior to the fall of the hegemony, by [Theorem 4](#), the theory predicts that there should be a small fraction but large number of periods where there are failed rebellions: lands that are lost to other societies but quickly regained. These failed rebellions may or may not involve zealots and need not take place when ξ_j is at its nadir. However, prior to the actual exit, ξ_j must be such that state power is at a nadir, and this means that resistance to rebellions is lower, so there should be more frequent and larger rebellions prior to the final fall of the society.

The exit path must have least resistance to leaving the basin. We know one such path in which the zealots simply grab one unit of land after another. This implies that the hegemon can lose land only to the zealots—since any other loss would incur greater resistance—and that it can never regain land—since then the zealots would have to regain it. This is exactly as in the discussion of transitions in the example of [Section 2](#).

Once the inner basin is breached we enter the outer range. We call this the *fall of the hegemony*. Once the hegemony loses resistance there will be several societies competing, each with an appreciable chance of success. We refer to this turbulent period before the basin of another hegemony is reached as the period of *warring states*. During this period, there will be many societies that may rise and fall, and swap land back and forth: it is a chaotic and turbulent phase. The exit path to another recurrent communicating class will then be concluded with a *rise* of a new hegemony. The rise of the new hegemony is in some respects opposite of the fall. Once a stable society has enough land that it has positive resistance to any opponents (implied if they have positive resistance to an opponent consisting entirely of zealots), least resistance implies it can only gain land and not lose it.

The entire period of transition we know ([Theorem 4](#)) to be short relative to the length of the hegemony and, of course, that must be true individually for each of the three phases. We would like to say something about the relative length of the three phases. The fall and the rise are both monotone: during the fall the hegemony only loses land; during the rise, the new hegemony only gains land. Hence we expect these phases to be relatively fast. By contrast, during the warring states land may swap back and forth many times before a victor is established. This leads us to suspect that the warring states period should last considerably longer than the fall and the rise. However, if we simply fix the number of units of land, no such result is possible: the warring states period may involve only a few units of land, while the fall and the rise could involve many units of land. In this case—since just a few units of land are involved—the back and forth during the warring states period does not much matter, since there is a good chance one society quickly gains the small number of units of land needed to establish a new hegemony.

Warring states with many units of land It is natural to think in terms of relatively small units of land, for example, the Alsace–Lorraine region, a relatively small area swapped back and forth between France and Germany four times in less than a century. To capture “small units” of land, we model the idea that the number of units of land L is large. In this case, we might expect that the length of the rise and fall—being monotone—will last an amount of time roughly proportional to L , while the warring state, which is more like a random walk, would last an amount of time more nearly proportional to L^2 .

To make this precise, define the share of land held by society j as $\theta_j = L_j/L$, and suppose that the conflict resolution and land gain resistance functions are continuous functions of these shares. We also suppose that there is a threshold $\bar{\theta} > \frac{1}{2}$ such that if no society has this fraction of land, then the land gain resistance of all active societies must be zero. Notice that this implies that any path from one hegemony to another must enter the warring states phase, since at some point the original hegemon falls below $\bar{\theta}L$ units of land for the first time and at this point the warring state phase is entered, since certainly no other society has that much land.²³

As we increase L , the number of states, which consists of all ways of dividing L units of land among a fixed number J of societies, grows very rapidly. Hence bounds on the expected length of direct routes that depend on the number of states are not going to be particularly useful. In [Appendix A](#), specific bounds for least resistance paths are given that do not depend on the number of states and a stronger version of this is given in [Appendix S.A](#). These bounds show that the least resistance paths corresponding to the rise and fall have an expected length proportional to L .

We also want to argue that the warring state phase lasts much longer than the fall or the rise. To do this we need an additional assumption: we need to know not just resistance during the warring states phase, but something about the actual transition probabilities. We assume that during the warring states phase, when no society has more than $\bar{\theta}L$ units of land each active society has the same probability $\beta < 1/J$ of losing or gaining a unit of land. Notice that at some point a new would-be hegemon must have $\frac{1}{2}L$ units of land. Hence take an initial condition in which the society with the most land is j and $L_{jt} = \frac{1}{2}L$ units of land. Then $L_{j\tau}$ is a random walk with β chance of increasing by one or decreasing by one at least until either $L_{j\tau} \geq \bar{\theta}L$ or $L_{j\tau} \leq (1 - \bar{\theta})L$.²⁴ In [Appendix S.B](#), we establish that the expected passage time is of order L^2 , which can be compared to the expected length of the rise and fall in [Appendix S.B](#) that are of order L . We put these results together in [Appendix S.C](#) to establish the following proposition.

PROPOSITION 2. *For any K there exists an $\bar{L}$ such that for all $L \geq \bar{L}$ there exists an $\bar{\epsilon}$ such that for all $\epsilon \leq \bar{\epsilon}$ the expected length of the warring states phase exceeds that of either the fall or the rise by K periods.*

Note the order of limits here: for larger L , we will generally have to choose smaller ϵ .

7.2 The fall of the last Qing dynasty in China

An interesting exercise is to compare the theoretical predictions of the transition to the fall of an actual hegemony. As a case study for which there is quite a bit of historical information, we take the fall of the final Qing dynasty in China and the subsequent rise of

²³Since we are assuming L is large, we are ignoring the rounding off needed due to the integer constraint.

²⁴Even if $L_{j\tau} \leq (1 - \bar{\theta})L$, no other society may have enough land to become a hegemon, but certainly no other society can have enough land to become a hegemon until this condition is satisfied, so we can use this condition to derive a lower bound on the expected length of the warring states period.

the communist hegemony.²⁵ The basic fact is that Chinese institutions that lasted from roughly the introduction of the Imperial Examination System in 605 CE until 1911 CE were swept away in less than a year. It is useful to begin the story about 1838, before the First Opium War. At that time, the Qing dynasty held a hegemony over China proper: the area bordered by the difficult terrain of Indochina in the Southeast, the Himalayan mountains in the South, the inhospitable deserts in the West, the Pacific Ocean in the East and the wasteland of Mongolia in the North. It also held a number of outlying areas not part of China proper: the Korean Peninsula, Indochina, and Taiwan. As these are not so easily defended, are not Chinese, and have only been part of the Chinese hegemony at certain times—and moreover, the current government claims only Taiwan among these territories—we do not count them as part of the hegemony.

Several independent sources of instability concurred to the fall of the hegemony. In the early 1800s China fell into a severe economic depression from which it did not recover prior to the fall of the hegemony. Outsiders, most notably the English, French, and Japanese, actively intervened in China, sometimes fighting for and other times against the Qing, but in any case certainly piling on pressure. Opium consumption, induced by the English to correct trade imbalances, increased as well.

From 1839 to 1910 there were a series of unsuccessful attempts to overthrow the Qing dynasty including local rebellions and acts of defiance by committed revolutionaries. During this time the outlying territories were lost: Korea became independent, Indochina was lost to the French, and Taiwan was lost to the Japanese, further weakening the hegemon. Roughly speaking, the state ξ_j became increasingly worse. However, each internal rebellion was successfully repressed, each war brought to an end, and in each case, the Qing hegemony over China proper—tax collecting authority, and control of local and global institutions—remained intact. There were institutional changes that took place during this period, some forced by the outsiders, and an attempt to placate the revolutionaries, such as the abolition of the imperial examination system in 1905. These can be viewed as shocks ξ_j that further weakened the state. Although it is hard to measure the relative frequency of failed rebellions before and after the economic weakness of the 19th century, in the earlier periods there seem not to have been such dramatic episodes as the Boxer rebellion and the less known Duggan revolt (which lasted for 15 years). As [Theorem 4](#) predicts, before the actual fall the state ξ_j is very bad, and there are many and probably increasing failed attempts at rebellion.

The actual fall of the Qing occurred in 1911 and as [Theorem 1](#) suggests, it was very quick. There were again a series of revolts; now, however, they succeeded. Also as the theory suggests, the length of the successful revolt (less than a year) is considerably shorter than the longest failed rebellions—the Boxer and Duggan rebellion lasting many years. The final successful revolt was coordinated by Sun Yat Sen. The groups carrying out the various revolts can reasonably be described as zealots: they share in common a dedication to overthrowing the Qing; they are willing to suffer severe risk and live under unpleasant circumstances to achieve that goal. Such behavior is power maximizing

²⁵There are of course many accounts of this period, and while they sometimes disagree on exactly who did what to whom and when, all agree on the basic facts we describe below. One readable account by a journalist is that of [Fenby \(2008\)](#).

but is not stable in the sense that no society has ever lasted very long based on the fanatical devotion of its members; neither was it the case in China. Hence the theoretical description of the fall of the hegemony is relatively accurate: zealots quickly capture the land and do so without a serious setback. In some cases, land is seized by other groups, but they quickly join Sun Yat Sen as the theory suggests. By the end of 1911, the Qing Emperor abdicated and Sun Yat Sen became the provisional President of China, which, however, no longer was hegemonic in any reasonable sense of the word.

Next is the period of warring states, both in theory and in fact. The theory says that there can be many competing societies, land may be lost and gained, and zealots may or may not play a role. Again, this is an accurate description of the situation in China between 1911 and 1946. Sun Yat Sen was quickly deposed by a less fanatical and more materialistic warlord Yuan Shikai, but until about 1927, and even after, there were many warlords in various parts of China who rise and fall, and many revolutions, some successful and others unsuccessful. There is also the Sino–Tibetan war and the Soviet invasion of Xinjiang during this period. Basically, the theory predicts chaos (in the nontechnical sense) and that is what we see. Beginning in about 1927, things settle down slightly with two relatively more powerful groups, the Nationalists and the Communists, fighting a civil war, but there remain many warlords who continue to rise and fall, at times forming alliances or professing allegiance to the two more significant groups. These two groups, unlike the earlier revolutionaries, appear to have coherent and potentially stable institutions. Then in 1936 the Japanese seize control of most of the country, an occupation that lasts until 1945. Notice that as the theory suggests, the length of this warring states period (35 years) is much longer than either the fall (less than 1 year) and the rise (about 3 years).

The final stage of a least resistance transition is the rise of the new hegemon. Again all transitions must have zero resistance, but now we are in the basin of the hegemony, so the least resistance path consists of the hegemony gaining territory—without losing any—until hegemony is again established. Notice that since, in this model, once the basin is left there are zero-resistance transitions to any particular hegemony breaching the threshold, the model makes no prediction about which hegemony eventually emerges; in particular, there is a nonnegligible probability that even a very weak hegemony emerges. In China, the threshold appears to be reached about 1946 when the Communists controlled about a quarter of the country and about a third of the population. They quickly overran the remaining areas held by the Nationalists, who retreated to Taiwan in 1949.

8. CONCLUSION

This paper is about events and combinations of events that are unlikely and that can be modeled as a finite Markov process, in particular how such a process moves from one relatively stable long-run state to another. Examples are transitions between different equilibria in a game or different political regimes. We show that these systems exhibit long periods of stability punctuated by brief episodes of change, and we give a detailed description of the probabilities and frequencies of these different outcomes. Within the

literature on “evolution of conventions,” we complement the results of [Kandori et al. \(1993\)](#), [Young \(1993\)](#), [Ellison \(2000\)](#), [Cui and Zhai \(2010\)](#), and [Hasker \(2014\)](#) on long-run dynamics in games. When applied to the context of social evolution, our theory has implications both for the societies we are likely to see and for the design of institutions: institutions that will persist for long periods of time must be robust against multiple failures, and it is these multiple failures that lead a society to fall.

It may be useful to look at smaller systems about which we have a great deal of information—also subject to small unlikely shocks, and subject to the same type of Markov analysis—to see what is involved; for example, commercial airlines, which crash relatively infrequently despite the vast number of flights and miles flown. As our theory predicts, when they do crash, it is typically due to multiple near simultaneous failures. To take a specific example, on November 24, 2001, en route from Berlin on approach to Zurich, Crossair Flight 3597 crashed near Bassersdorf, Switzerland, killing 24 of the 33 people on board. According to the flight investigation, seven independent unfortunate events occurred on that occasion.²⁶ These multiple failures seem typical of commercial aviation crashes. Each individual failure is unlikely, but none terribly so. What is highly unlikely is that all occur in combination. In general, airplanes are designed with a high degree of redundancy to provide insulation against failure of one or even several components: multiple pilots, multiple navigation systems, multiple engines, multiple independent hydraulic systems, and so forth. So it is with human societies. For example, penal codes and the legal systems have a high degree of redundancy (appeals procedures) to prevent the punishment of the innocent. Societies that survive for long periods of time must be well cushioned against even multiple failures. For example, the fall of the Roman Empire has been attributed to many factors: religious ferment, the plague, corruption, the forced migration of hostile outsiders, economic recession, and so forth. Despite the effort of historians to establish each as “the” cause of the fall, as is the case with Flight 3597, all of these things happened, and while each is uncommon, none is particularly unlikely, and the Roman Empire had suffered through each of these, often in combination, many times before. What is unique about the fall is that all these things occurred at once. When a system or society is well designed, it takes a perfect storm—everything going wrong at once—to bring it down. But, as this paper shows, it is the least unlikely combination of things—the least resistance direct route—that will typically lead, for good or ill, to abrupt and sudden change.

APPENDIX A: DIRECT ROUTES

Our goal is to establish probability and expectations bounds on subsets $A \subseteq A_{xBW} \neq \emptyset$. Define $t(A) = \min\{t(a) \mid a \in A, r(a) = r(A)\}$ to be the minimum number of transitions of

²⁶See [AAIB \(2002\)](#): (i) the pilot had a bad record of following procedures during landing and was inadequately trained, but was allowed nevertheless to transport passengers; (ii) the flight was behind schedule and, consequently, the pilot was in a hurry to land; (iii) due to noise regulations, the plane was diverted to a less safe runway; (iv) the runway had inadequate instrumentation, and the airport parameters and protocols for landing on the runway were inadequate; (v) the range of hills the plane crashed into was not marked on the chart; (vi) the pilot put the plane into an overly steep descent and descended too low without proper visual contact with the ground; (vii) the pilot did not monitor the proper instruments during the attempted landing.

any least resistance path in the set A . This appendix is devoted to proving the following main result on direct paths that characterizes their probability and length. [Theorem 1](#) in the text follows directly.

THEOREM 11. *For $A \subseteq A_{xBW}$ nonempty and $t \geq 0$, there are bounds $G(t), H(t) > 0$ nondecreasing in t such that $C^{t(A)} \epsilon^{r(A)} \leq P_\epsilon(A|x) \leq G(t(A)) \epsilon^{r(A)}$ and $E[t(a)|x, A] \leq H(t(A))$.*

The bounds G and H will be specifically computed.

First we establish that $r(A) = \min_{a \in A} r(a)$ exists. Then getting a lower bound on $P_\epsilon(A|x)$ is relatively easy: it is bounded below by the probability of a path $a \in A$ with resistance $r(a)$, which is to say, it is of order $\epsilon^{r(A)}$. The main goal is to establish a similar upper bound. The problem is that A can easily contain infinitely many paths with resistance $r(A)$ as well as paths of greater resistance. However, there are only finitely many paths of any given length, so if there are infinitely many paths, most of them must be very long. The idea is that since paths in A must avoid the comprehensive set W , they are not likely to be very long since there are zero-resistance routes to W . To make this precise, we construct a finite set of template paths of relatively low resistance and show that all the paths in A can be constructed by adding loops to the template paths. We then show that the probability of all paths constructed from a given template is bounded by the probability of the template times a constant that does not depend on ϵ . This same method also yields bounds on the expected length of the direct routes.

Loop cutting and a lower bound

In general paths $a \in A$ contain loops, and since an analysis of loops forms a key part of the analysis, we begin by introducing the notion of loop cutting. The idea is to construct templates for a from which a can be reconstructed by adding loops. If $a = (z_0, z_1, \dots, z_t)$, we say that a' is a *loop cut* of a at $z_\tau = z_{\tau'}$ for $t > \tau' > \tau \geq 0$ if $a' = (z_0, z_1, \dots, z_\tau, z_{\tau'+1}, \dots, z_t)$, that is, if the loop at z_τ has been cut out.²⁷ Note the obvious fact that $r(a') \leq r(a)$.

DEFINITION 7. A map m from the set of all paths to itself is a *loop-cutting algorithm* if there is a sequence $a_1, a_2, \dots, a_M$ with $a_1 = a$, $a_M = m(a)$ and a_{j+1} a loop cut of a_j for $j = 1, 2, \dots, M-1$.

Note that $r(m(a)) \leq r(a)$. The path $m(a)$ is a template for a from which a can be reconstructed by adding loops. A loop-cutting algorithm is *maximal* if $m(m(a)) = m(a)$.

We can establish the existence of $\min_{a \in A} r(a)$ using the *zero-cut* algorithm defined as follows. For any $a = (z_0, z_1, \dots, z_t)$, if there is no loop of zero-resistance, stop; otherwise, cut the first and shortest loop of zero resistance and repeat. This is obviously maximal. Note that $m(a)$ is a *no-zero-loop* path in the sense that it contains no zero-resistance

²⁷Note that there cannot be a loop cut that begins with the final element z_t . Indeed $z_t \in B$; when the path gets there it stops.

loops and that $r(m(a)) = r(a)$. Our first step is to give a bound on the length of no-zero-loop paths. Let Z_A to be the set of non-end points touched by paths in A , that is, the set of z_τ such that there is some $(x, z_1, z_2, \dots, z_\tau, \dots, z_t) \in A$ with $\tau < t$ —leading and trailing elements not counted. Let N_A be the number of elements in Z_A , and let N_{AB} be the maximum of N_A and the number of elements in the target B . These are both bounded above by N , but may be much smaller: using computations based only on A allows our results to be extended from a finite state space to a countable state space. Let $\underline{r}(A)$ be the smallest finite nonzero resistance of any transition in any path in A .

LEMMA 1. *If $a \in A$ is a no-zero-loop path,²⁸ then $t(a) \leq N_A r(a) / \underline{r}(A)$.*

PROOF. Observe that since nonzero-resistance transitions have resistance at least $\underline{r}(A)$, there are at most $r(a) / \underline{r}(A)$ such transitions in a , and the remaining transitions must have zero resistance. Since there are no zero-resistance loops, the number of zero-resistance transitions between each positive resistance transition is at most N_A . $\square$

We can now apply the zero-cut algorithm to prove the following basic fact.

LEMMA 2. *The function $r(A) = \min_{a \in A} r(a)$ is well defined.*

PROOF. Fix $a \in A$. Recall that, by assumption, a has positive probability for $\epsilon > 0$, so that $r(a) < \infty$. Let m be the zero-cut algorithm. Consider that for any $a' \in A$ with $r(a') \leq r(a)$, we have $r(a') = r(m(a'))$. Let $\bar{A}$ be the set of all finite resistance paths $(x, z_1, z_2, \dots, z_\tau, \dots, z_t)$ with $z_t \in B$ and $z_\tau \in Z_A$. Since $m(a') \in \bar{A}$ then by Lemma 1 $t(m(a')) \leq N_{\bar{A}} r(a') / \underline{r}(\bar{A})$. But there are only finitely many paths of length t that begin at x and take values in Z_A and end in B , so only finitely many possible values of $r(a') \leq r(a)$. Hence $\min_{a \in A} r(a)$ exists. $\square$

Having established that $r(A)$ is well defined, we can easily establish a lower bound on $P_\epsilon(A|x)$.

LEMMA 3. *We have $P_\epsilon(A|x) \geq C^{t(A)} \epsilon^{r(A)}$.*

PROOF. Let $a \in A$ satisfy $r(a) = r(A)$ and $t(a) = t(A)$. Then

$$P_\epsilon(A|x) \geq P_\epsilon(a|x) = \prod_{\tau=1}^{t(a)} P_\epsilon(z_\tau | z_{\tau-1}) \geq \prod_{\tau=1}^{t(a)} C \epsilon^{r(z_\tau, z_{\tau-1})} = C^{t(a)} \epsilon^{r(a)}. \quad \square$$

More about loops To establish an upper bound, we need further results about loops. First, we give a useful refinement on Lemma 1 using the *all-cut* algorithm, which is defined as follows. For any $a = (z_0, z_2, \dots, z_t)$, if there is no loop, stop; otherwise cut the first and shortest loop, and repeat. This is obviously maximal.

²⁸Note that this results for any set of paths A , not just direct routes.

LEMMA 4. *If $a \in A$ is a no-zero-loop path, then*

$$t(a) \leq 2N_A \left(1 + \frac{r(a) - r(A)}{\underline{r}(A)} \right).$$

PROOF. Let m be the all-cut algorithm. Observe that to get from $m(a)$ to a , we must add loops containing at least $t(a) - t(m(a))$ elements. Suppose that there are fewer than $(t(a) - t(m(a)))/(2N_A)$ loops that cannot be subdivided into two non-overlapping sub-loops. Then one of these loops must have length at least $2N_A$. But the first N_A of the loop must contain a loop and so must the second N_A so that the loop can be divided into two non-overlapping sub-loops. By assumption, none of the loops have zero resistance, so each has resistance at least $\underline{r}(A)$; hence

$$r(a) \geq r(m(a)) + (t(a) - t(m(a))) \frac{\underline{r}(A)}{2N_A}$$

and the result follows from $r(m(a)) \geq r(A)$ and $t(m(a)) \leq N_A$. $\square$

Next we consider a loop-cutting algorithm that produces templates with resistance no smaller than $r(A)$. We say that m *preserves r* if $r(a) \geq r$ implies $r(m(a)) \geq r$. One such algorithm is the r -preserving algorithm. For any $a = (x, z_1, \dots, z_t)$, if no loop can be cut without reducing the resistance of a below r , stop; otherwise, cut the first and shortest such loop and repeat. Observe that for an r -preserving algorithm, the image $m(A)$ consists of no-zero-loop paths and is maximal since removing any loop would necessarily reduce the resistance below r . The key property of this algorithm is that it produces templates with resistance not too much bigger than r and of bounded length (by Lemma 1), and in particular that means there are finitely many templates.

DEFINITION 8. We take $\bar{r}(A)$ to be the greatest (finite) resistance of any transition in any path in A , and

$$\bar{t}(A) \equiv 2N_A \left(2 + \frac{r(A) - r(A_{xBW})}{\underline{r}(A)} \right).$$

LEMMA 5. *For $a \in A$ and the $r(A)$ -preserving loop-cut algorithm, $r(m(a)) \leq r(A) + N_A \bar{r}(A)$ and $t(m(a)) \leq \bar{t}(A)$, and $m(A)$ is finite with at most $N_{AB}^{\bar{t}(A)}$ elements.*

PROOF. Observe first that if m is the $r(A)$ -preserving loop-cut algorithm if $r(m(a)) > N_A \bar{r}(A)$, then $m(a)$ must have a loop of resistance greater than zero and, hence, must have a loop of resistance no greater than $N_A \bar{r}(A)$. If $r(m(a)) > r(A) + N_A \bar{r}(A)$, removing such a loop leaves resistance greater than or equal to $r(A)$, contradicting the fact that the $r(A)$ -preserving algorithm can leave no such loop.

Now let m be the $r(A)$ -preserving algorithm and suppose that $r(m(a)) > r(A_{xBW})$. Remove the shortest loop from $m(a)$ to get a' so that $r(m(a)) \geq r(A) > r(a') \geq r(A_{xBW})$. Since a' has no zero-resistance loops, by Lemma 4,

$$t(a') \leq 2N_A \left(1 + \frac{r(a') - r(A_{xBW})}{\underline{r}(A)} \right).$$

Also, since the a' is $m(a)$ with the shortest loop—so less than N_A in length—removed, we must have $t(m(a)) \leq t(a') + N_A$. Hence

$$t(m(a)) \leq 2N_A \left(2 + \frac{r(a') - r(A_{xBW})}{\underline{r}(A)} \right) \leq 2N_A \left(2 + \frac{r(A) - r(A_{xBW})}{\underline{r}(A)} \right) = \bar{t}(A).$$

If $r(m(a)) = r(A)$, then m removes all the loops, so we have the bound $t(m(a)) \leq N_A$ so that certainly $t(m(a)) \leq \bar{t}(a)$ and that bound holds in all cases.

Finally, since $a \in m(A)$ are constructed of elements from $Z_A \cup B$ and have length at most $\bar{t}(A)$, there are at most $N_{AB}^{\bar{t}(A)}$ elements. $\square$

We also can give a bound in terms of $t(a)$ as follows.

LEMMA 6. *We have*

$$\bar{t}(A) \leq 2N_A \left(2 + \frac{t(A)\bar{r}(A)}{\underline{r}(A)} \right).$$

PROOF. Clearly $r(A) \geq 0$. Moreover, since there is a path $a \in A$ with length $t(a) = t(A)$, such a path can have resistance at most $t(A)\bar{r}(A)$, so $r(A)$ cannot be greater than this. $\square$

Upper bounds

Once the loops have been removed to create templates, we must put them back in to compute probabilities. It is convenient at this point to work with loops with the first element removed, so that a loop at z_τ is a sequence of the form $\emptyset, (z_\tau)$ or $(\zeta_\tau, \dots, \zeta_{\tau+k}, z_\tau)$, where $\zeta_\tau \in Z_A$.²⁹ If we have a transition $(z_\tau, z_{\tau+1})$ and ℓ_τ is a loop at z_τ , then the corresponding path is $a = (z_\tau, \ell_\tau, z_{\tau+1})$, with number of transitions $t(a)$ equal to the number of elements of ℓ_τ plus 1.

For any path $a = (x, z_1, z_2, \dots, z_t)$ and $0 \leq \tau \leq t$, let $a[\tau] = (x, z_1, z_2, \dots, z_\tau)$ be the corresponding τ truncation. Then for any $a = (z_0, z_1, z_2, \dots, z_t) \in A$, $\tau < t$, and path a' , consider the path $(a[\tau], a') = (z_0, z_1, z_2, \dots, z_\tau, a')$ that starts along a and deviates to a' at z_τ . Define $t_F(a, \tau)$ to be the least length of any path (z_τ, a') that has zero resistance and $(a[\tau], a') \notin A$ (so that the deviation does not reach B). Notice that $t_F(a, \tau) \leq N - 1$: since W is comprehensive, there is a path (z_τ, a') of zero resistance with no loops that ends in W , and such a path can have at most N elements, hence at most $N - 1$ transitions.

DEFINITION 9. The *failure time* $t_F(A) = \max_{a \in A, \tau < t(a)} t_F(a, \tau)$.

Hence $t_F(A) \leq N - 1$, but it may be much smaller: in one of the examples in the text, $t_F(A) = 1$ regardless of N .

To establish upper bounds on the probability of A and on the expected length of its paths, we start by establishing the fact that, given that W is comprehensive, long loops

²⁹Note that we must include $\emptyset$ because when we compute probabilities of transitions from z to y along a set of loops, we must also include the probability that we go directly from z to y without a loop, that is, along a null loop.

are not very likely. Let $\lfloor R \rfloor$ be the largest integer not greater than R . We can now give the following bound on the probability of long loops.

LEMMA 7. *Suppose $z \in Z_A$ and that $L_A(z)$ is a set of loops at z . Then for $t \geq 0$, we have $P_\epsilon(t((z, \ell, y)) = t + 1, \ell \in L_A(z)|z) \leq P_\epsilon(y|z)(1 - C^{t_F(A)})^{\lfloor t/t_F(A) \rfloor}$.*

PROOF. If $t = 0$, the result is immediate since in fact the left and right sides are equal. Since (z, ℓ, y) always ends with the transition (z, y) , we have

$$P_\epsilon(t((z, \ell, y)) = t + 1, \ell \in L_A(z)|z) \leq P_\epsilon(y|z)$$

for all t and in particular for $t < t_F(A)$, which is the desired bound in that case. If $t \geq t_F(A)$, we may define L_k to be the set of paths (z, ℓ) with $\ell \in L_A(z)$ truncated at $kt_F(A)$ for $1 \leq k \leq \lfloor t/t_F(A) \rfloor$ and we have

$$P_\epsilon(t((z, \ell, y)) = t + 1, \ell \in L_A(z)|z) \leq P_\epsilon(L_{\lfloor t/t_F(A) \rfloor}|z)P_\epsilon(y|z).$$

Moreover, $L_{\lfloor t/t_F(A) \rfloor} \subseteq L_k$ for $1 \leq k \leq \lfloor t/t_F(A) \rfloor$, so it suffices to prove recursively that $P_\epsilon(L_k|z) \leq (1 - C^{t_F(A)})^k$ for $1 \leq k \leq \lfloor t/t_F(A) \rfloor$.

First we take $k = 1$. Observe that starting at z , there is a zero-resistance path of length no longer than $t_F(A)$ that reaches a point y that is contained in no loop. Since the probability of each zero-resistance transition is at least C , there is at most a probability $1 - C^{t_F(A)}$ of remaining in the set L_1 .

Now we suppose the result is true for k and prove it for $k + 1$. Each loop in L_{k+1} has the form (a_1, z', a_2) , where $(a_1, z') \in L_k$ and $t(a_2) = t_F(A)$. Let $L^+(a_1, z')$ be the set $\{a_2 \mid (a_1, z', a_2) \in L_{k+1}\}$. Then

$$\begin{aligned} P_\epsilon(L_{k+1}|x) &= \sum_{(a_1, z') \in L_k} \sum_{a_2 \in L^+(a_1, z')} P_\epsilon((a_1, z')|x)P_\epsilon(a_2|z') \\ &= \sum_{(a_1, z') \in L_k} P_\epsilon((a_1, z')|x) \sum_{a_2 \in L^+(a_1, z')} P_\epsilon(a_2|z') \\ &= \sum_{(a_1, z') \in L_k} P_\epsilon((a_1, z')|x)P_\epsilon(L^+(a_1, z')|z'). \end{aligned}$$

Moreover, since again there is a zero-resistance path starting at z' of length no longer than $t_F(A)$ that reaches a point y' that is contained in no loop, we have $P_\epsilon(L^+(a_1, z')|z') \leq 1 - C^{t_F(A)}$. Hence

$$\begin{aligned} P_\epsilon(L_{k+1}|x) &\leq \sum_{(a_1, z') \in L_k} P_\epsilon((a_1, yz')|x)(1 - C^{t_F(A)}) \\ &= P_\epsilon(L_k|x)(1 - C^{t_F(A)}) \\ &\leq (1 - C^{t_F(A)})^{k+1} \end{aligned}$$

by the inductive hypothesis. □

We are now ready to reverse the loop-cutting procedure by adding loops to templates to construct the paths in A . Opposite to a loop-cutting algorithm is the idea of a *loop insertion set*. Let $L_A(z_\tau)$ be a set of loops at z_τ . We suppose we are given an $a = (z_0, z_1, \dots, z_t) \in A$. A loop set $\overline{m^{-1}}(a)$ is defined by a sequence of sets of loops $L_\tau \subseteq L_A(z_\tau)$, $\tau = 0, 1, \dots, t-1$, and consists of all paths of the form $(z_0, \ell_0, z_1, \ell_1, \dots, z_{t-1}, \ell_{t-1}, z_t)$ such that $\ell_\tau \in L_\tau$. If m is a loop-cutting algorithm and $a \in m(A)$ if $\overline{m^{-1}}(a) \supseteq A \cap m^{-1}(a)$, we say that $\overline{m^{-1}}(a)$ covers m, a .

We now define

$$S_k(\overline{m^{-1}}(a)|x) \\ \equiv \sum_{t_0=0}^{\infty} \sum_{t_1=0}^{\infty} \cdots \sum_{t_{t(a)-1}=0}^{\infty} \left[\left(\sum_{s=0}^{t(a)-1} t_s \right)^k \right] \prod_{\tau=0}^{t(a)-1} P_\epsilon(t((z_\tau, \ell, z_{\tau+1})) = t_\tau + 1, \ell \in L_\tau | z_\tau).$$

The significance of these numbers is given by the following lemma.

LEMMA 8. If $\overline{m^{-1}}(a)$ covers m, a , then

$$\sum_{a' \in \overline{m^{-1}}(a)} P_\epsilon(a'|x) \leq S_0(\overline{m^{-1}}(a)|x) \\ \sum_{a' \in \overline{m^{-1}}(a)} t(a') P_\epsilon(a'|x) \leq S_1(\overline{m^{-1}}(a)|x)$$

and both hold with equality if every $a' \in \overline{m^{-1}}(a)$ has a unique representation of the form $a' = (z_0, \ell_0, z_1, \ell_1, \dots, z_t)$, where $\ell_\tau \in L_\tau$.

PROOF. By definition every $a' \in \overline{m^{-1}}(a)$ has a representation of the form $a' = (z_0, \ell_0, z_1, \ell_1, \dots, z_t)$, where $\ell_\tau \in L_\tau$.³⁰ Hence for any nonnegative function $f(a')$, we have

$$\sum_{a' \in \overline{m^{-1}}(a)} f(a') P_\epsilon(a'|x) \\ \leq \sum_{\ell_0 \in L_0} \sum_{\ell_1 \in L_1} \cdots \sum_{\ell_{t(a)-1} \in L_{t(a)-1}} f((z_0, \ell_0, z_1, \ell_1, \dots, z_t)) \prod_{\tau=0}^{t(a)-1} P_\epsilon(\ell_\tau | z_\tau)$$

with equality if the representation is unique; that is, if each a' has a unique representation, then it appears exactly once in the sum. Rearranging the sum by adding over the length of the loops then gives the desired result. $\square$

We can now compute the desired upper bounds.

³⁰A simple example shows that there can be more than one representation. Suppose $(z_0, z_1, z_2) \in m(A)$. If $\ell_0 = (z_1, z_0)$, $\ell_1 = \emptyset$, then $a' = (z_0, \ell_0, z_1, \ell_1, z_2) = (z_0, z_1, z_0, z_1, z_2)$. If $\ell_0 = \emptyset$, $\ell_1 = (z_0, z_1)$, then $(z_0, \ell_0, z_1, \ell_1, z_2) = (z_0, z_1, z_0, z_1, z_2) = a'$.

LEMMA 9. *We have*

$$S_0(\overline{m^{-1}}(a)|x) \leq P_\epsilon(a|x)[t_F(A)/C^{t_F(A)}]^{t(a)}$$

and

$$S_1(\overline{m^{-1}}(a))/S_0(\overline{m^{-1}}(a)) \leq 3t(a)t_F(A)^2/C^{2t_F(A)}.$$

PROOF. From Lemma 7 we have

$$\begin{aligned} S_0(\overline{m^{-1}}(a)|x) &= \sum_{t_0=0}^{\infty} \sum_{t_1=0}^{\infty} \cdots \sum_{t_{t(a)-1}=0}^{\infty} \prod_{\tau=0}^{t(a)-1} P_\epsilon(t((z_\tau, \ell, z_{\tau+1})) = t_\tau + 1, \ell \in L_\tau | z_\tau) \\ &= \prod_{\tau=0}^{t(a)-1} \sum_{t=0}^{\infty} P_\epsilon(t((z_\tau, \ell, z_{\tau+1})) = t + 1, \ell \in L_\tau | z_\tau) \\ &\leq \prod_{\tau=0}^{t(a)-1} \sum_{t=0}^{\infty} P_\epsilon(z_{\tau+1} | z_\tau) (1 - C^{t_F(A)})^{\lfloor t/t_F(A) \rfloor} \\ &= \left(\prod_{\tau=0}^{t(a)-1} P_\epsilon(z_{\tau+1} | z_\tau) \right) \prod_{\tau=0}^{t(a)-1} \sum_{t=0}^{\infty} (1 - C^{t_F(A)})^{\lfloor t/t_F(A) \rfloor} \\ &= P_\epsilon(a|x)[t_F(A)/C^{t_F(A)}]^{t(a)}. \end{aligned}$$

Next to simplify notation, set

$$P_\tau(t_\tau) \equiv P_\epsilon(t((z_\tau, \ell, z_{\tau+1})) = t_\tau + 1, \ell \in L_\tau | z_\tau).$$

Then

$$\begin{aligned} S_1(\overline{m^{-1}}(a)|x) &= \sum_{t_0=0}^{\infty} \sum_{t_1=0}^{\infty} \cdots \sum_{t_{t(a)-1}=0}^{\infty} \left(\sum_{s=0}^{t(a)-1} t_s \right) \prod_{\tau=0}^{t(a)-1} P_\tau(t_\tau) \\ &= \sum_{s=0}^{t(a)-1} \sum_{t_0=0}^{\infty} \sum_{t_1=0}^{\infty} \cdots \sum_{t_{t(a)-1}=0}^{\infty} t_s \prod_{\tau=0}^{t(a)-1} P_\tau(t_\tau) \\ &= \sum_{s=0}^{t(a)-1} \left(\frac{\sum_{t=0}^{\infty} t_s P_s(t)}{\sum_{t=0}^{\infty} P_s(t)} \right) \prod_{\tau=0}^{t(a)-1} \sum_{t=0}^{\infty} P_\tau(t) \\ &= \sum_{s=0}^{t(a)-1} \left(\frac{\sum_{t=0}^{\infty} t_s P_s(t)}{\sum_{t=0}^{\infty} P_s(t)} \right) S_0(\overline{m^{-1}}(a)|x). \end{aligned}$$

Moreover,

$$\begin{aligned} \sum_{t=0}^{\infty} P_s(t) &\geq P_s(1) = P_\epsilon(t((z_s, \ell, z_{s+1})) = 1, \ell \in L_s | z_s) \\ &= P_\epsilon(z_{s+1} | z_s), \end{aligned}$$

and applying again [Lemma 7](#) and using a summation formula proven in [Appendix S.A](#), we have

$$\begin{aligned} S_1(\overline{m^{-1}}(a)|x)/S_0(\overline{m^{-1}}(a)|x) &\leq \sum_{s=0}^{t(a)-1} \left(\frac{\sum_{t_s=0}^{\infty} t_s P_{\epsilon}(z_{s+1}|z_s)(1 - C^{t_F(A)})^{\lfloor t_s/t_F(A) \rfloor}}{P_{\epsilon}(z_{s+1}|z_s)} \right) \\ &= t(a) \sum_{t_s=0}^{\infty} t_s (1 - C^{t_F(A)})^{\lfloor t_s/t_F(A) \rfloor} \\ &\leq t(a) 3t_F(A)^2 / C^{2t_F(A)}, \end{aligned}$$

which is the desired bound. $\square$

We can now establish the upper bounds stated in [Theorem 11](#).

THEOREM 12. *If $A \subseteq A_{xBW}$ is nonempty, then $P_{\epsilon}(A|x) \leq [N_{AB} D t_F(A) / C^{t_F(A)}]^{\tilde{t}(A)} \epsilon^{r(A)}$.*

PROOF. Take m to be the $r(A)$ -preserving loop-cut algorithm and, for $a \in m(A)$, take $\overline{m^{-1}}(a)$ to be defined by the sequence $L_{\tau} = L_A(z_{\tau})$. Let $T(A) = \max_{a \in m(A)} t(a) \leq \tilde{t}(A)$. Since $P_{\epsilon}(\overline{m^{-1}}(a)|x) \leq S_0(\overline{m^{-1}}(a)|x)$, we may apply [Lemma 9](#) and use the fact that $\overline{m^{-1}}(a)$ covers m, a to conclude $P_{\epsilon}(A|x) \leq \#m(A) \epsilon^{r(A)} [D t_F(A) / C^{t_F(A)}]^{T(A)}$. Moreover, by [Lemma 5](#), we have $\#m(A) \leq N_{AB}^{\tilde{t}(A)}$, giving the desired result. $\square$

THEOREM 13. *If $A \subseteq A_{xBW}$ is nonempty, then*

$$E[t(a)|x, A] \leq \tilde{t}(A) \left[\frac{3t_F(A)^2}{C^{2t_F(A)}} \right] \frac{[N_{AB} D t_F(A) / C^{t_F(A)}]^{\tilde{t}(A)}}{C^{t(A)}},$$

and if A are the least resistance paths and $\tilde{t}(A)$ is the longest least resistance path containing no loops, then

$$E[t(a)|x, A] \leq \tilde{t}(A) \left[\frac{3t_F(A)^2}{C^{2t_F(A)}} \right].$$

PROOF. In all cases, by [Lemma 9](#), we have for any loop-cutting algorithm that preserves $r(A)$,

$$\begin{aligned} E[t(a)|x, A] &\leq \frac{\sum_{a \in m(A)} S_1(\overline{m^{-1}}(a)|x)}{P_{\epsilon}(A|x)} \\ &= \sum_{a \in m(A)} \frac{S_1(\overline{m^{-1}}(a)|x)}{S_0(\overline{m^{-1}}(a))} \frac{S_0(\overline{m^{-1}}(a))}{P_{\epsilon}(A|x)} \\ &\leq \frac{3t_F(A)^2}{C^{2t_F(A)}} \sum_{a \in m(A)} \frac{t(a) S_0(\overline{m^{-1}}(a))}{P_{\epsilon}(A|x)}. \end{aligned}$$

In general, we can take m to be the $r(A)$ -preserving loop-cut algorithm and, for $a \in m(A)$, take $\overline{m^{-1}}(a)$ to be defined by the sequence $L_\tau = L_A(z_\tau)$. Then we observe that there are at most $N_{AB}^{\tilde{t}(A)}$ templates and we apply Theorems 3 and 12 to get

$$E[t(a)|x, A] \leq \frac{3t_F(A)^2}{C^{2t_F(A)}} N_{AB}^{\tilde{t}(A)} \frac{\tilde{t}(A) D^{\tilde{t}(A)} \epsilon^{r(A)} [t_F(A)/C^{t_F(A)}]^{\tilde{t}(A)}}{C^{t(A)} \epsilon^{r(A)}},$$

giving the first result.

Now suppose that A are the least resistance paths. We can now take m to be the all-cut algorithm, which, since no least resistance path can have a positive resistance loop in it, is the same as the zero-cut algorithm when applied to A . For $a \in m(A)$, we now take $\overline{m^{-1}}(a)$ to be defined by the sequence L_τ of zero-resistance loops in $L_A(z_\tau)$ that do not contain z_s for $s < \tau$.³¹

First we observe that $\overline{m^{-1}}(a) = m^{-1}(a) \cap A$. Starting at z_τ , the all-cut algorithm cuts the longest loop ending at z_τ ; hence z_τ cannot subsequently appear in the template. Moreover, the loops added back are exactly the ones that were cut. Second we observe that for a given template $(z_1, z_2, \dots, z_t) \in m(A)$, and two states $a' = (z_0, \ell'_0, z_1, \ell'_1, \dots, z_t)$ and $a'' = (z_0, \ell''_0, z_1, \ell''_1, \dots, z_t) \in \overline{m^{-1}}(a)$, then $a' = a''$ only if $\ell'_0, \ell'_1, \dots, \ell'_{t-1} = \ell''_0, \ell''_1, \dots, \ell''_{t-1}$. Hence $S_0(\overline{m^{-1}}(a)) = P_\epsilon(m^{-1}(a) \cap A)$. So

$$\begin{aligned} E[t(a)|x, A] &\leq \frac{3t_F(A)^2}{C^{2t_F(A)}} \sum_{a \in m(A)} \frac{t(a) P_\epsilon(m^{-1}(a) \cap A)}{P_\epsilon(A|x)} \\ &= \frac{3t_F(A)^2}{C^{2t_F(A)}} \sum_{a \in m(A)} \frac{t(a) P_\epsilon(m^{-1}(a) \cap A)}{\sum_{a \in m(A)} P_\epsilon(m^{-1}(a) \cap A)} \\ &\leq \tilde{t}(A) \frac{3t_F(A)^2}{C^{2t_F(A)}}, \end{aligned}$$

giving the second result. □

APPENDIX B: QUASI-DIRECT ROUTES

We fix a start point $x \in \Omega_x$, a target set B , and a quasi-comprehensive forbidden set $W \supseteq B$, and we let Q_{xBW} be the corresponding nonempty set of quasi-direct routes. Recall that the paths $a \in Q_{xBW}$ can be decomposed as $a_1, a_2, \dots, a_{n(a)}, a^+$, where the $a_i \in A_{xxW}$ are loops from x to x that do not touch x or W in between and $a^+ \in A_{xBW}$ is the exit path, a direct route from x to B that does not touch W or x in between.

THEOREM 3 (Repeated). *Let $A = \{a \in Q_{xBW} \mid \rho(a) = \rho(Q_{xBW})\}$ denote the least peak resistance paths in Q_{xBW} . Then $\lim_{\epsilon \rightarrow 0} (P_\epsilon(A|x)/P_\epsilon(Q_{xBW} \setminus A|x)) = \infty$.*

PROOF. Fix $\rho = \rho(Q_{xBW})$. By Theorem 2, the set A consists of the set of paths of the form $a_1, a_2, \dots, a_n, a^+$, where $a_i \in A_{xxW}$ and $a^+ \in A_{xBW}$, and $r(a_i) \leq \rho$, $r(a^+) = \rho$. Define the set of paths $A_{>\rho}$ as those having the form $a_1, a_2, \dots, a_n, a_{n+1}$, where $a_i \in A_{xxW}$

³¹Notice that the construction in the example of footnote 30 is ruled out in the present case.

and $a_{n+1} \in A_{xxW} \cup A_{xBW}$, and $r(a_i) \leq \rho$, $r(a_{n+1}) > \rho$. We claim that $P_\epsilon(A_{>\rho}|x) \geq P_\epsilon(Q_{xBW} \setminus A|x)$ so that it will suffice to prove that $\lim_{\epsilon \rightarrow 0} (P_\epsilon(A|x)/P_\epsilon(A_{>\rho}|x)) = \infty$. To prove this claim, observe that if $a \in Q_{xBW} \setminus A$, then the first part of the path must necessarily lie in $A_{>\rho}$ so the event $Q_{xBW} \setminus A$ implies the event $A_{>\rho}$.

Now let A_{xxW}^ρ , A_{xBW}^ρ and $A_{xxW}^{>\rho}$, $A_{xBW}^{>\rho}$ denote the subsets of resistance exactly equal to ρ and strictly bigger than ρ , respectively. We compute

$$\begin{aligned} \frac{P_\epsilon(A|x)}{P_\epsilon(A_{>\rho}|x)} &= \frac{\sum_{n=0}^{\infty} P_\epsilon^n(A_{xxW}^\rho|x) P_\epsilon(A_{xBW}^\rho|x)}{\sum_{n=0}^{\infty} P_\epsilon^n(A_{xxW}^\rho|x) P_\epsilon(A_{xxW}^{>\rho} \cup A_{xBW}^{>\rho}|x)} \\ &= \frac{P_\epsilon(A_{xBW}^\rho|x)}{P_\epsilon(A_{xxW}^{>\rho} \cup A_{xBW}^{>\rho}|x)} \geq \frac{P_\epsilon(A_{xBW}^\rho|x)}{P_\epsilon(A_{xxW}^{>\rho}|x) + P_\epsilon(A_{xBW}^{>\rho}|x)} \end{aligned}$$

and the result now follows directly from [Theorem 1](#) on the probability of direct paths. $\square$

When $B = W$, the decomposition also makes it easy to do computations since the loops a_i are independent and identically distributed random variables. Specifically, for $f: A_{xxW} \rightarrow \mathfrak{R}$ and $a \in A$, define $F(a) \equiv \sum_{i=1}^{n(a)} f(a_i)$. Then for any function $g(n)$ of the number of loops we have the following lemma.

LEMMA 10. *If $B = W$, then $E(Fg|x, Q_{xBW}) = E(f|x, A^0)E(ng|x, Q_{xBW})$ and $P_\epsilon(n|x, Q_{xBW})$ is geometric with $E(n|x, Q_{xBW}) = (1/P_\epsilon(A_{xBW}|x)) - 1$.*

PROOF. Since $B = W$, we have $P_\epsilon(A_{xxW}|x) + P_\epsilon(A_{xBW}|x) = 1$, while A_{xxW} and A_{xBW} are disjoint. Then abbreviating $Q = Q_{xBW}$, we get

$$\begin{aligned} E(Fg|x, Q) &= E\left[\sum_{i=1}^n f(a_i)g|x, Q\right] \\ &= E\left[\sum_{i=1}^n E[f(a_i)g|x, Q, n]|x, Q\right] = E\left[\sum_{i=1}^n gE[f(a_i)|x, Q, n]|x, Q\right]. \end{aligned}$$

The event (Q, n) is exactly the event $a_i \in A_{xxW}$ for $i = 1, 2, \dots, n$ and $a_{n+1} \in A_{xBW}$, and conditional on x these are independent events. Hence $E(f(a_i)|x, Q, n) = E(f|x, A_{xxW})$. We conclude that

$$\begin{aligned} E(Fg|x, Q) &= E\left[\sum_{i=1}^n gE(f|x, A^0)|x, Q\right] \\ &= E(f|x, A^0)E\left[\sum_{i=1}^n g|x, Q\right] = E(f|x, A^0)E(ng|x, Q). \end{aligned}$$

This is the first result. Also since $P_\epsilon(A_{xxW}|x) + P_\epsilon(A_{xBW}|x) = 1$, it follows that n is geometrically distributed with success probability $P_\epsilon(A_{xBW}|x)$, which gives the stated expected value. $\square$

Recall that $M(a, A)$ is the number of loops of a that lie in $A \subseteq A_{xxW}$; that is, if $f : A_{xxW} \rightarrow \mathbb{R}$ is the indicator of A (taking the value 1 if $a^0 \in A$ and 0 otherwise), then $M(a, A)$ is the aggregate F defined by $F(a, A) \equiv \sum_{i=1}^{n(a)} f(a_i)$. Also, recall that $t^-(a)$ is the amount of time along a spent outside of $\Omega(x)$. Let Z_{xW} be the subset of Z that is reachable by finite resistance paths that start at x and touch W at most once, and let N_{xW} be the number of elements in Z_{xW} . This is bounded above by N , and sets A of direct routes that we consider satisfy $Z_A \subseteq Z_{xW}$ so that $N_{AB} \leq N_{xW}$. Define

$$G_1 \equiv [N_{xW}^2 D / C^{N_{xW}}]^{N_{xW}}, \quad H_1 = \frac{6N_{xW}^3 G_1}{C^{3N_{xW}}}.$$

THEOREM 4 (Repeated). *If $B = W$, then*

$$\begin{aligned} \epsilon^{-r(A_{xBW})} / G_1 &\leq E_\epsilon[t(a)|x, Q_{xBW}], E_\epsilon[t(a^-)|x, Q_{xBW}] \\ &\leq H_1(1 + C^{-2N_{xW}} \epsilon^{-r(A_{xBW})}) \epsilon^{-r(A_{xBW})} \end{aligned}$$

while

$$\lim_{\epsilon \rightarrow 0} E_\epsilon \left[\frac{t^-(a)}{t(a^-)} \middle| x, Q_{xBW} \right] = 0.$$

For $A \subseteq A_{xxW}$,

$$G_1^{-1} C^{2N_{xW}} \epsilon^{r(A) - r(A_{xBW})} \leq E_\epsilon[M(a, A)|x, Q_{xBW}] \leq G_1 C^{-2N_{xW}} \epsilon^{r(A) - r(A_{xBW})},$$

and if $r(A) < r(A_{xBW})$, then for all $k \geq 0$,

$$\lim_{\epsilon \rightarrow 0} P_\epsilon[M(a, A) > k|x, Q_{xBW}] = 1.$$

PROOF. Observe that for $A \in \{A_{xxW}, A_{xBW}\}$, $t_F(A) \leq N_{xW}$ and $\bar{t}(A) = 4N_A \leq 4N_{xW}$ so that the bound from [Theorem 12](#) is in turn bounded by G_1 and that from [Theorem 13](#) is bounded by H_1 .

From [Lemma 10](#), $E_\epsilon[t(a^-)|x, Q_{xBW}] \leq E_\epsilon[t|x, A_{xxW}] / P_\epsilon(A_{xBW}|x)$. Moreover, recall that $t(A)$ is the number of transitions in the shortest of the least resistance paths in A ; since A_{xxW} contains all zero-resistance loops $r(A_{xxW}) = 0$, the shortest of these loops is no longer than N_{xW} , so $t(A_{xxW}) \leq N_{xW}$. Analogously, A_{xBW} contains templates without loops; hence $t(A_{xBW}) \leq N_{xW}$. Hence from [Theorems 12 and 13](#), we have $1 \leq E_\epsilon[t|x, A_{xBW}] \leq H_1$, $1 \leq E_\epsilon[t|x, A_{xxW}] \leq H_1$ and $C^{N_{xW}} \epsilon^{r(A_{xBW})} \leq P_\epsilon(A_{xBW}|x) \leq G_1 \epsilon^{r(A_{xBW})}$. This gives the stated bounds on $E_\epsilon[t|x, Q_{xBW}]$.

Now set $\tau^- = t^-(a)$ and $\tau = t(a^-)$. For the next part, observe that $E_\epsilon[\tau^- / \tau | x, Q_{xBW}] \leq E_\epsilon[\tau^- / n | x, Q_{xBW}] = E_\epsilon[t^-(a) | x, A_{xxW}]$. Now split A_{xxW} into two disjoint sets $\mathcal{A}_0^0$ of paths of zero resistance and $\mathcal{A}_r^0$ of positive resistance, where r is the least positive resistance in A_{xxW} . Then $E_\epsilon[t^- | x, A_{xxW}] = E_\epsilon[t^- | x, \mathcal{A}_0^0] P_\epsilon[\mathcal{A}_0^0 | x, A_{xxW}] + E_\epsilon[t^- | x, \mathcal{A}_r^0] P_\epsilon[\mathcal{A}_r^0 | x, A_{xxW}]$. However, $E_\epsilon[t^- | x, \mathcal{A}_0^0] = 0$ by definition, while by [Theorems 12 and 13](#)

$$E_\epsilon[t^- | x, \mathcal{A}_r^0] P_\epsilon[\mathcal{A}_r^0 | x, A_{xxW}] \leq [H_1 G_1 / C^{t(\mathcal{A}_r^0)}] \epsilon^r \rightarrow 0,$$

which implies the result.

Next, given $A \subseteq A_{xxW}$ and f the indicator of A , [Lemma 10](#) gives

$$\begin{aligned} E_\epsilon[M(a, A)|x, Q_{xBW}] &= E_\epsilon[F|x, Q_{xBW}] = E_\epsilon(f|x, A_{xxW})/P_\epsilon(A_{xBW}|x) \\ &= P_\epsilon(A|x, A_{xxW})/P_\epsilon(A_{xBW}|x). \end{aligned}$$

Applying [Theorem 12](#) (and recalling that $r(A_{xxW}) = 0$) then gives the stated bounds on $M(a, A)$.

Last, since $P_\epsilon(n|x, Q_{xBW})$ is geometric, we have

$$P_\epsilon(n(a) \geq n|x, Q_{xBW}) = (1 - P_\epsilon(A_{xBW}|x))^n$$

and

$$\begin{aligned} P_\epsilon[M(a, A) > k|x, Q_{xBW}, n(a) = n] &= 1 - \sum_{i=0}^k \binom{n}{i} P_\epsilon(A|x)^i (1 - P_\epsilon(A|x))^{n-i} \\ &\geq 1 - (k+1)n^k(1 - P_\epsilon(A|x))^n. \end{aligned}$$

By hypothesis we can choose $r(A) < r < r(A_{xBW})$. Take $n = \epsilon^{-r}$. Then by [Theorem 12](#), we have $P_\epsilon(n(a) \geq n|x, Q_{xBW}) \geq (1 - G_1\epsilon^{r(A_{xBW})})^{\epsilon^{-r}}$ and $P_\epsilon[M(a, A) > k|x, Q_{xBW}, n(a) = n] \geq 1 - (k+1)\epsilon^{-rk}(1 - G_1\epsilon^{r(A)})^{\epsilon^{-r}}$. Taking the log of the first expression and using de l'Hospital's rule gives, as $\epsilon \rightarrow 0$,

$$\lim_{\epsilon \rightarrow 0} \frac{\log(1 - G_1\epsilon^{r(A_{xBW})})}{\epsilon^r} = \lim_{\epsilon \rightarrow 0} -\frac{G_1r(A_{xBW})\epsilon^{r(A_{xBW})-1}}{r\epsilon^{r-1}(1 - G_1\epsilon^{r(A_{xBW})})} = \lim_{\epsilon \rightarrow 0} -\frac{G_1r(A_{xBW})}{r}\epsilon^{r(A_{xBW})-r} = 0$$

so that $P_\epsilon(n(a) \geq n|x, Q_{xBW}) \rightarrow 1$. Next take the log of $\epsilon^{-rk}(1 - G_1\epsilon^{r(A)})^{\epsilon^{-r}}$ to find

$$-rk \log \epsilon + \epsilon^{-r} \log(1 - G_1\epsilon^{r(A)}) = \left[1 - \frac{rk \log \epsilon}{\epsilon^{-r} \log(1 - G_1\epsilon^{r(A)})} \right] \epsilon^{-r} \log(1 - G_1\epsilon^{r(A)}).$$

Then by de l'Hospital's rule,

$$\begin{aligned} \lim_{\epsilon \rightarrow 0} \frac{rk \log \epsilon}{\epsilon^{-r} \log(1 - G_1\epsilon^{r(A)})} &= \lim_{\epsilon \rightarrow 0} \frac{rk}{-\frac{r \log(1 - G_1\epsilon^{r(A)})}{\epsilon^r} - \frac{G_1r(A)\epsilon^{r(A)-r}}{\log(1 - G_1\epsilon^{r(A)})}} \\ &\leq -\lim_{\epsilon \rightarrow 0} \frac{rk \log(1 - G_1\epsilon^{r(A)})}{G_1r(A)\epsilon^{r(A)-r}} = 0 \end{aligned}$$

and

$$\lim_{\epsilon \rightarrow 0} \frac{\log(1 - G_1\epsilon^{r(A)})}{\epsilon^r} = \lim_{\epsilon \rightarrow 0} -\frac{G_1r(A)\epsilon^{r(A)-r}}{r(1 - G_1\epsilon^{r(A)})} = -\infty$$

so $-rk \log(\epsilon) + \epsilon^{-r} \log(1 - G_1\epsilon^{r(A)}) \rightarrow -\infty$ and $\epsilon^{-rk}(1 - G_1\epsilon^{r(A)})^{\epsilon^{-r}} \rightarrow 0$. Hence $P_\epsilon[M(a, A) > k|x, Q_{xBW}, n(a) = n] \rightarrow 1$. $\square$

Recalling that $A_{xxW}(t)$ are the loops that spend at least t consecutive periods outside of Ω_x , we now prove the corollary.

COROLLARY 6 (Final part of [Theorem 5](#)). *If there is a path $a^0 \in A_{xxW}$ that contains a zero-resistance loop not touching $\Omega(x)$ with $0 < r(a^0) < r(A_{xBW})$, then for any $k > 0$,*

$$\lim_{\epsilon \rightarrow 0} P_\epsilon(M(a, A_{xxW}(kt(a^+))) > k|x, Q_{xBW}) = 1.$$

PROOF. Fix a $\delta > 0$. By [Theorem 11](#) and Chebychev's inequality we may choose K such that $P_\epsilon(t(a^+) > K) < 1 - \delta$. Given a^0 as described, we can insert zero-resistance loops to get a path $a^K \in A_{xxW}(K)$ with resistance $r(a^K) < r(A_{xBW})$. Hence we may apply [Theorem 4](#) to conclude that for all sufficiently small ϵ ,

$$P_\epsilon[M(a, A_{xxW}(K)) > K|x, Q_{xBW}] < 1 - \delta.$$

Since, conditional on x , the loops and exit path are independent, we then have

$$P_\epsilon(M(a, A_{xxW}(kt(a^+))) > K|x, Q_{xBW}) < (1 - \delta)^2,$$

which implies the desired result. $\square$

APPENDIX C: ERGODIC PROBABILITIES AND BOUNDS

[THEOREM 7](#) (Repeated). *If $y \in \Omega_x$, then*

$$\lim_{\epsilon \rightarrow 0} \frac{\mu_\epsilon(x)}{\mu_\epsilon(y)} = \frac{\bar{\mu}_0(x)}{\bar{\mu}_0(y)}.$$

PROOF. Partition the matrix P_ϵ with rows corresponding to source states and columns corresponding to target states into P_ϵ^{ij} , where $i, j = 1$ corresponds to Ω_x and $i, j = 2$ corresponds to $\Omega \setminus \Omega_x$. In particular P^{11} is square, the size of Ω_x . Correspondingly, let e^i be the column vectors of 1s with dimension corresponding to $i = 1, 2$. Define the row vector $\bar{\mu}_\epsilon(z) = \mu_\epsilon(z) / \sum_{y \in \Omega(x)} \mu_\epsilon(y)$ and partition this vector conformally. Since $\bar{\mu}_\epsilon$ is normalized to 1 on Ω_x and $\bar{\mu}_0$ is strictly positive, it suffices to prove that as $\epsilon \rightarrow 0$, every limit point $\bar{\mu}_\epsilon^1$ is equal to $\bar{\mu}_0^1$, where we include the superscript to emphasize that we are dealing only with the invariant distribution on Ω_x . The invariance condition is $\bar{\mu}_\epsilon^1 = \bar{\mu}_\epsilon^1 P_\epsilon^{11} + \bar{\mu}_\epsilon^2 P_\epsilon^{21}$. Multiplying this on the right by e^1 , we get $1 = \bar{\mu}_\epsilon^1 P_\epsilon^{11} e^1 + \bar{\mu}_\epsilon^2 P_\epsilon^{21} e^1$, while the fact that P_ϵ is a Markov kernel means that $P_\epsilon^{11} e^1 + P_\epsilon^{12} e^2 = e^1$ or $P_\epsilon^{11} e^1 = e^1 - P_\epsilon^{12} e^2$. Substituting, we see that $1 = \bar{\mu}_\epsilon^1 (e^1 - P_\epsilon^{12} e^2) + \bar{\mu}_\epsilon^2 P_\epsilon^{21} e^1 = 1 - \bar{\mu}_\epsilon^1 P_\epsilon^{12} e^2 + \bar{\mu}_\epsilon^2 P_\epsilon^{21} e^1$ so that $\bar{\mu}_\epsilon^2 P_\epsilon^{21} e^1 = \bar{\mu}_\epsilon^1 P_\epsilon^{12} e^2$, which says roughly that the steady state flow into Ω_x must equal the steady state flow out. Now $P_\epsilon^{12} \rightarrow 0$ as $\epsilon \rightarrow 0$ since these are the probabilities of leaving the recurrent communicating set Ω_x ; hence $\bar{\mu}_\epsilon^2 P_\epsilon^{21} e^1 \rightarrow 0$. But $\bar{\mu}_\epsilon^2 P_\epsilon^{21}$ is a nonnegative vector, so $\bar{\mu}_\epsilon^2 P_\epsilon^{21} e^1 \rightarrow 0$ is possible only if $\bar{\mu}_\epsilon^2 P_\epsilon^{21} \rightarrow 0$. Then in the invariance condition $\bar{\mu}_\epsilon^1 = \bar{\mu}_\epsilon^1 P_\epsilon^{11} + \bar{\mu}_\epsilon^2 P_\epsilon^{21}$, as $P_\epsilon^{11} \rightarrow P_0^{11}$ if $\bar{\mu}_0^1$ is a limit point of $\bar{\mu}_\epsilon^1$, it must satisfy the limiting condition that $\bar{\mu}_0^1 = \bar{\mu}_0^1 P_0^{11}$. However, as Ω_x is recurrent, communicating this equation has only one solution $\bar{\mu}_0^1$, so we conclude that, in fact, $\bar{\mu}_\epsilon^1 \rightarrow \bar{\mu}_0^1$. $\square$

Recall that $r(\Omega_x) = \min_{y \in \Omega} r(\Omega_x, \Omega_y)$. Also let $\bar{r}$, $\underline{r}$ be the largest resistance of any transition and smallest positive resistance, and set $\bar{G} \equiv G_1((N-1)(1 + (\bar{r}/\underline{r})))$.

THEOREM 8 (Repeated). *Allowing that Ω_x may be empty, if $A = A_{xyW}$ are the direct routes from x to y with forbidden set $W = \{x\} \cup \{y\} \cup (\Omega \setminus \Omega_x)$, then $\mu_\epsilon(y) \geq \mu_\epsilon(x)C^N\epsilon^{r(A)}$. If $x \in \Omega_x$ and there is a zero-resistance path from y to x , then also $\mu_\epsilon(y) \leq \mu_\epsilon(x)C^{-N}\overline{G}^2\epsilon^{\min\{r(A), r(\Omega_x)\}}$.*

PROOF. We use the standard fact about Markov ergodic probabilities as used for example by Ellison (2000): If we let $N_\epsilon(y, x|x)$ be the expected number of times y occurs before x starting at x , then $\mu_\epsilon(y) = \mu_\epsilon(x)N_\epsilon(y, x, |x)$.

The lower bound is immediate: Since with probability $P_\epsilon(A)$, we have y hit once without returning to x , we have from Theorem 11 $\mu_\epsilon(y) = \mu_\epsilon(x)N_\epsilon(y, x|x) \geq \mu_\epsilon(x)P_\epsilon(A) \geq \mu_\epsilon(x)C\epsilon^{r(A)}$.

Next we suppose that y has zero resistance for getting to $x \in \Omega_x$. We use the reverse condition $\mu_\epsilon(x) = \mu_\epsilon(y)N_\epsilon(x, y, |y)$, so we must find a lower bound on $N_\epsilon(x, y|y)$. Let $A_1 = A_{yx(x \cup y \cup (\Omega \setminus \Omega_x))}$. Observe that $N_\epsilon(x, y|y) \geq P_\epsilon(A_1)N_\epsilon(x, y|x)$. Since there is a zero-resistance path from y to x , we have from Theorem 11 the bound $P_\epsilon(A_1) \geq C^N$, so $N_\epsilon(x, y|y) \geq C^N N_\epsilon(x, y|x)$.

Now define sets $B = \{y\} \cup (\Omega \setminus \Omega_x)$ and $A_2 = A_{xB(x \cup B)}$. Then $N_\epsilon(x, y|x) \geq N_\epsilon(x, B|x)$. Since starting at x the events B and $\sim B = A_2$ are mutually exclusive independent events, $N_\epsilon(x, B|x) = 1/P_\epsilon(\sim B) = 1/P_\epsilon(A_2)$. Since from Theorem 11 $P_\epsilon(A_2) \leq \epsilon^{r(A_2)}$, we get $N_\epsilon(x, y|y) \geq C^N \epsilon^{-r(A_2)}/\overline{G}$, or $\mu_\epsilon(y) \leq \mu_\epsilon(x)C^{-N}\overline{G}^2\epsilon^{r(A_2)}$.

Finally, the event A_2 is contained in the event $A_{xy(x \cup y \cup (\Omega \setminus \Omega_x))} \cup A_{x(\Omega \setminus \Omega_x)(x \cup (\Omega \setminus \Omega_x))}$. Hence

$$r(A_2) = \min\{r(A_{xy(x \cup y \cup (\Omega \setminus \Omega_x))}), r(A_{x(\Omega \setminus \Omega_x)(x \cup (\Omega \setminus \Omega_x))})\}.$$

However $r(A_{x(\Omega \setminus \Omega_x)(x \cup (\Omega \setminus \Omega_x))}) = r(\Omega_x)$ and $r(A_{xy(x \cup y \cup (\Omega \setminus \Omega_x))}) = r(A)$, which gives the desired upper bound. $\square$

REFERENCES

- Acemoglu, Daron and James A. Robinson (2001), "A theory of political transitions." *American Economic Review*, 91, 938–963. [91, 108]
- Aircraft Accident Investigation Bureau (AAIB) (2002), "Final report concerning the accident to the aircraft AVRO 146-RJ100, HB-IXM, operated by Crossair under flight number CRX 3597, on 24 November 2001 near Bassersdorf/ZH." Report 1793. [116]
- Binmore, Kenneth and Larry Samuelson (1997), "Muddling through: Noisy equilibrium selection." *Journal of Economic Theory*, 74, 235–265. [91]
- Blume, Lawrence E. (2003), "How noise matters." *Games and Economic Behavior*, 44, 251–271. [91]
- Cui, Zhiwei and Jian Zhai (2010), "Escape dynamics and equilibria selection by iterative cycle decomposition." *Journal of Mathematical Economics*, 46, 1015–1029. [90, 105, 116]
- Ellison, Glenn (1993), "Learning, local interaction, and coordination." *Econometrica*, 61, 1047–1071. [91]

Ellison, Glenn (2000), “Basins of attraction, long-run stochastic stability, and the speed of step-by-step by evolution.” *Review of Economic Studies*, 67, 17–45. [90, 91, 93, 96, 97, 99, 100, 104, 105, 106, 107, 110, 116, 130]

Ellison, Glenn, Drew Fudenberg, and Lorens A. Imhof (2014), “Fast convergence in evolutionary models: A Lyapunov approach.” Unpublished paper. [89]

Fenby, Jonathan (2008), *Modern China: The Fall and Rise of a Great Power, 1850 to the Present*. Penguin Press, London. [114]

Foster, Dean P. and H. Peyton Young (1990), “Stochastic evolutionary game dynamics.” *Theoretical Population Biology*, 38, 219–232. [89]

Hasker, Kevin (2014), “The emergent seed: A representation theorem for models of stochastic evolution and two formulas for waiting time.” Working paper, Bilkent University. [90, 105, 107, 116]

Kandori, Michihiro, George J. Mailath, and Rafael Rob (1993), “Learning, mutation, and long run equilibria in games.” *Econometrica*, 61, 29–56. [89, 91, 116]

Kreindler, Gabriel E. and H. Peyton Young (2012), “Rapid innovation diffusion with local interaction.” Economics Discussion Paper Series, University of Oxford. [89]

Kreindler, Gabriel E. and H. Peyton Young (2013), “Fast convergence in evolutionary equilibrium selection.” *Games and Economic Behavior*, 80, 39–67. [89]

Levine, David K. and Salvatore Modica (2013a), “Anti-malthus: Conflict and the evolution of societies.” *Research in Economics*, 67, 289–306. [102, 108]

Levine, David K. and Salvatore Modica (2013b), “Conflict, evolution, hegemony and the power of the state.” Working Paper 19221, NBER. [90, 91, 108]

Sabourian, Hamid and Wei-Torng Juang (2012), “A folk theorem on equilibrium selection.” Working paper, Cambridge University. [89]

Young, H. Peyton (1993), “The evolution of conventions.” *Econometrica*, 61, 57–84. [89, 90, 116]